

Rings, circles, and null-models for point pattern analysis in ecology

Thorsten Wiegand and Kirk A. Moloney

Wiegand, T. and Moloney, K. A. 2004. Rings, circles, and null-models for point pattern analysis in ecology. – *Oikos* 104: 209–229.

A large number of methods for the analysis of point pattern data have been developed in a wide range of scientific fields. First-order statistics describe large-scale variation in the intensity of points in a study region, whereas second-order characteristics are summary statistics of all point-to-point distances in a mapped area and offer the potential for detecting both different types and scales of patterns. Second-order analysis based on Ripley's K-function is increasingly used in ecology to characterize spatial patterns and to develop hypothesis on underlying processes; however, the full range of available methods has seldomly been applied by ecologists. The aim of this paper is to provide guidance to ecologists with limited experience in second-order analysis to help in the choice of appropriate methods and to point to practical difficulties and pitfalls. We review (1) methods for analytical and numerical implementation of two complementary second-order statistics, Ripley's K and the O-ring statistic, (2) methods for edge correction, (3) methods to account for first-order effects (i.e. heterogeneity) of univariate patterns, and (4) a variety of useful standard and non-standard null models for univariate and bivariate patterns. For illustrative purpose, we analyze examples that deal with non-homogeneous univariate point patterns. We demonstrate that large-scale heterogeneity of a point-pattern biases Ripley's K-function at smaller scales. This bias is difficult to detect without explicitly testing for homogeneity, but we show that it can be removed when applying methods that account for first-order effects. We synthesize our review in a number of step-by-step recommendations that guide the reader through the selection of appropriate methods and we provide a software program that implements most of the methods reviewed and developed here.

T. Wiegand, Dept of Ecological Modelling, UFZ-Centre for Environmental Research, PF 500136, DE-04301 Leipzig, Germany (tow@oesa.ufz.de). – K. A. Moloney, Dept of Botany, 143 Bessey Hall, Iowa State Univ., Ames, Iowa 50011-1020, USA.

Over the last decade, there has been an increasing interest in the study of spatial patterns in ecology (Turner 1989, Levin 1992, Grimm et al. 1996, Gustafson 1998, Dale 1999, Liebhold and Gurevitch 2002). Ecologists study spatial pattern to infer the existence of underlying processes (Perry et al. 2002). For example, spatial patterns of plants may result from different

processes and forces such as seed dispersal, intraspecific competition, interspecific competition, disturbance, herbivory, or environmental heterogeneity (Sturner et al. 1986, Kenkel 1988, Barot et al. 1999, Dale 1999, Jeltsch et al. 1999, Klaas et al. 2000), which may operate at different spatial scales. Analysis of the resulting spatial structure can indicate the existence of underlying

Accepted 21 August 2003

Copyright © OIKOS 2004
ISSN 0030-1299

processes, e.g. identification of regularity within a spatial pattern may indicate competition. However, care is required in inferring causation because many different processes may generate the same spatial pattern.

An important example of a spatial pattern is a point pattern, a data-set consisting of a series of mapped point locations in some study region. A number of methods have been used for quantifying the characteristics of point patterns (Ripley 1981, Diggle 1983, Upton and Fingleton 1985, Stoyan and Stoyan 1994, Bailey and Gatrell 1995, Dale 1999, Dale et al. 2002). First-order statistics describe the intensity λ of a point pattern, and large-scale variation in the intensity λ of the points in the study region. In contrast, second-order statistics are based on the distribution of distances of pairs of points (Ripley 1981) and they describe the small-scale spatial correlation structure of the point pattern. Some of the second-order statistics, such as the commonly used Ripley's K-function or the pair-correlation function g , use the information on all inter-point distances (Ripley 1976, 1977, 1981, Diggle 1983, Bailey and Gatrell 1995) and provide more information on the scale of the pattern than do statistics that use nearest neighbor distances only (i.e. Diggle's nearest neighbor functions G or F ; Diggle 1983, Barot et al. 1999). Ripley's K function and the pair-correlation function describe the characteristics of the point pattern over a range of distance scales, and can therefore detect mixed patterns (e.g. dispersion at smaller distances and aggregation at larger distances). This is an important property because virtually all ecological processes are scale dependent and their characteristics may change across scales (Levin 1992, Wiens et al. 1993, Gustafson 1998). Over the last several years, methods based on Ripley's K-function have undergone a rapid development (Haase 1995, Podani and Czárán 1997, Coomes et al. 1999, Goreaud and Pélissier 1999, Dale and Powell 2001, Haase 2001, Pélissier and Goreaud 2001, Dungan et al. 2002, Freeman and Ford 2002, Perry et al. 2002) and are now being widely used, especially in plant ecology (Haase et al. 1996, 1997, Ward et al. 1996, Martens et al. 1997, Hanus et al. 1998, Larsen and Bliss 1998, Pancer-Koteja et al. 1998, Wiegand et al. 1998, 2000, Barot et al. 1999, Chen and Bradshaw 1999, Coomes et al. 1999, Jeltsch et al. 1999, Mast and Veblen 1999, Kuusinen and Penttinen 1999, Bossdorf et al. 2000, Camarero et al. 2000, Condit et al. 2000, Grau 2000, He and Duncan 2000, Kuuluvainen and Rouvinen 2000, Lookingbill and Zavala 2000, Haase 2001, Fehmi and Bartolome 2001). Some examples can be found in other fields of ecology as well (O'Driscoll 1998, Wiegand et al. 1999, Klaas et al. 2000, Schooley and Wiens 2001, Revilla and Palomares 2002).

The function $K(r)$ is the expected number of points in a circle of radius r centered at an arbitrary point (which is not counted), divided by the intensity λ of the pattern. The alternative pair correlation function $g(r)$, which

arises if the circles of Ripley's K-function are replaced by rings (Ripley 1981, Stoyan and Stoyan 1994, Stoyan and Penttinen 2000, Dale et al. 2002), gives the expected number of points at distance r from an arbitrary point, divided by the intensity λ of the pattern. Of special interest is to determine whether a pattern is random, clumped, or regular. Significance is usually evaluated by comparing the observed data with Monte Carlo envelopes from the analysis of multiple simulations of a null model. The common null model is complete spatial randomness (CSR), but other null models may be appropriate depending on the biological question asked. Points in a point-pattern may contain information in addition to position, often referred to as marks (e.g. a species identifier, a life stage identifier, or whether the individual survived or died), and many biological questions concern the relationship between points with different marks (e.g. facilitation or competition among adult trees and seedlings, or shrubs and grasses). Bivariate extensions of Ripley's K and the pair-correlation function provide appropriate methods to address such questions.

Practical difficulties and pitfalls in univariate analysis arise due to edge effects (i.e. sample circles or rings fall partly outside the study region and cannot be evaluated without bias) or if the pattern is not homogenous (i.e. the intensity λ of the pattern is not approximately constant in the study region). Any spatial dependence that is indicated by the estimated K function of a heterogeneous pattern could be due more to first-order effects rather than to interaction between the points themselves. In this case a null model that acknowledges the overall first-order heterogeneity has to be adopted to examine possible second-order effects. Alternatively, one may examine homogeneous sub-regions of the heterogeneous pattern. The latter requires methods to delineate homogeneous sub-regions that may be in general arbitrarily shaped, and methods of edge correction for arbitrarily shaped study regions. The analysis of bivariate point patterns is more complicated than that of univariate patterns because various other null models in addition to CSR become possible. The appropriate null model for bivariate analysis must be selected carefully based on the biological hypothesis to be tested.

In this article, we review current methods in point pattern analysis based on second-order statistics and address practical difficulties and pitfalls of this technique. More specifically, we suggest the use of the O-ring statistic as useful complement to the commonly used Ripley's K-function, we review methods for (1) edge correction, (2) analytical and numerical implementation of second-order statistics, (3) delineating homogeneous sub-regions, and we thoroughly discuss null models for univariate and bivariate point patterns. For illustrative purpose, we show several examples that deal with non-homogeneous univariate patterns, and as synthesis of

our article we compile a set of step-by-step recommendations that guide through the explorative process of second-order statistics. Additionally we provide our own software that enables ecologists to use most methods reviewed here. It can be requested from the first author.

Methods

Ripley's K-function and the O-ring statistic

For a homogeneous and isotropic point pattern, the second-order characteristics depend only on distance r , but not on the direction or the location of points. An appropriate geometry is therefore to adopt circular shapes, such as the circles of Ripley's K-function, as a basis for the spatial statistics. Curiously, the alternative approach of using rings (or annulus) instead of circles, i.e. the pair-correlation function $g(r)$ or the O-ring statistic $O(r) = \lambda g(r)$ (Wiegand et al. 1999), has rarely been used in ecology (but see Galiano 1982, Wiegand et al. 2000, Condit et al. 2000, Revilla and Palomares 2002). Using rings instead of circles (Fig. 1) has the advantage that one can isolate specific distance classes, whereas the cumulative K-function confounds effects at larger distances with effects at shorter distances (Getis and Franklin 1987, Penttinen et al. 1992, Condit et al. 2000). Note that the K-function and the O-ring statistic respond to slightly different biological questions. The accumulative K-function can detect aggregation or dispersion up to a given distance r and is therefore appropriate if the process in question (e.g. the negative effect of competition) may work only up to a certain distance, whereas the O-ring statistic can detect aggrega-

tion or dispersion at a given distance r . The O-ring statistic has the additional advantage that it is a probability density function (or a conditioned probability spectrum, Galiano 1982) with the interpretation of a neighborhood density, which is more intuitive than an accumulative measure (Stoyan and Penttinen 2000). We therefore argue that the toolbox of second-order spatial analysis should include not only the cumulative K-function, but also the complementary O-statistic.

Ripley's K-function

The bivariate K-function $K_{12}(r)$ is defined as the expected number of points of pattern 2 within a given distance r of an arbitrary point of pattern 1, divided by the intensity λ_2 of points of pattern 2:

$$\lambda_2 K_{12}r = E[\#(\text{points of pattern 2} \leq r \text{ from an arbitrary point of pattern 1})] \quad (1)$$

where $\#$ means "the number of", and $E[\]$ is the expectation operator. Under independence of the two point patterns, $K_{12}(r) = \pi r^2$, without regard to the individual univariate point patterns. It can be difficult to interpret $K_{12}(r)$ visually. Therefore, a square root transformation of $K(r)$, called L-function (Besag 1977), is used instead:

$$L_{12}(r) = \left(\sqrt{\frac{K_{12}(r)}{\pi}} - r \right) \quad (2)$$

This transformation removes the scale dependence of $K_{12}(r)$ for independent patterns and stabilizes the variance (Ripley 1981). Values of $L_{12}(r) > 0$ indicate that there are on average more points of pattern 2 within distance r of points of pattern 1 as one would expect under independence, thus indicating attraction between the two patterns up to distance r . Similarly, values of $L_{12}(r) < 0$ indicate repulsion between the two patterns up to distance r . The estimated L-function $\hat{L}_{12}(r)$ is calculated for a sequence of distances r and the results of $\hat{L}_{12}(r)$ are then plotted against distance.

Theoretically, distribution theory could be used in determining confidence envelopes for null models of point-patterns. However, this approach quickly becomes analytically intractable if edge effects for irregularly shaped study regions are considered, or if null models other than CSR are considered. Therefore, the more practical alternative is to use Monte Carlo simulations of a realization of the stochastic process underlying the specific null model in constructing confidence envelopes around the null model (Upton and Fingleton 1985, Bailey and Gatrell 1995). Each simulation generates an $\hat{L}_{12}(r)$ function, and approximate $n/(n+1) \times 100\%$ confidence envelopes are calculated from the highest and lowest values of $\hat{L}_{12}(r)$ taken from n simulations of the null model. For example, a 95% confidence envelope requires $n = 19$ simulations (Bailey and Gatrell 1995,

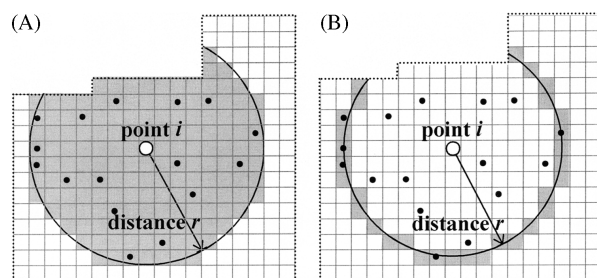


Fig. 1. Numerical implementation of the L-function and the O-ring statistic for an irregularly shaped study region encircled by the dashed line. Points of pattern 2 are represented by closed circles, the focal point i of pattern 1 as open circle. Note that we approximate circles and rings with the underlying grid structure. (A) For numerical implementation of Ripley's bivariate L-function we count the number of points of pattern 2 inside the part of the circles around point i of pattern 1 which falls inside the study region (i.e. the gray shaded area), and the number of cells within this area. (B) For implementation of the bivariate O-ring statistic we count the number of points of pattern 2 inside the part of the ring around point i of pattern 1 which falls inside the study region (i.e. the gray shaded area), and the number of cells within this area.

Haase 1995, Martens et al. 1997). A more accurate approach is to use the 5th-lowest and 5th highest $\hat{L}_{12}(r)$. In this case, 99 randomizations provide 5% confidence envelopes (Stoyan and Stoyan 1994, Wiegand et al. 2000). If $\hat{L}_{12}(r)$ has some part outside of that envelope, it is judged to be a significant departure from the null model.

The univariate K-function $K(r)$ is calculated in a manner analogous to the bivariate K function by setting pattern 1 equal to pattern 2. In this case the focal points of the circles are not counted. For a homogeneous Poisson process (complete spatial randomness CSR), $K(r) = \pi r^2$ and $L(r) = 0$. $L(r) > 0$ indicates aggregation of the pattern up to distance r , while $L(r) < 0$ indicates regularity of the pattern up to distance r .

The O-ring statistic

The mark-correlation function $g_{12}(r)$ is the analogue of Ripley's $K_{12}(r)$ when replacing the circles of radius r by rings with radius r , and the O-ring statistic $O_{12}(r) = \lambda_2 g_{12}(r)$ gives the expected number of points of pattern 2 at distance r from an arbitrary point of pattern 1 (Fig. 1B):

$$O_{12}(r) = \lambda_2 g_{12}(r) = E[\#(\text{points of pattern 2 at distance } r \text{ from an arbitrary point of pattern 1})] \quad (3)$$

The mark-correlation function $g_{12}(r)$ is related to Ripley's K-function (Ripley 1981, Stoyan and Stoyan 1994):

$$g_{12}(r) = \frac{dK_{12}(r)}{dr} / (2\pi r) \quad (4)$$

We obtain $O_{12}(r) = \lambda_2$ for independent patterns, $O_{12}(r) < \lambda_2$ for repulsion, whereas $O_{12}(r) > \lambda_2$ for attraction.

In practice, the calculation of the O-ring statistic involves a technical decision on the width of the rings. Clearly, the use of rings that are too narrow will produce jagged plots as not enough points will fall into the different distance classes. This problem does not occur for the accumulative K-functions. On the other hand, the O-ring statistic will lose the advantage that it can isolate specific distance classes if the rings are too wide.

Again, the univariate O-ring statistic $O(r)$ is calculated by setting pattern 2 equal to pattern 1. For CSR, $O(r) = \lambda$, $O(r) > \lambda$ indicates aggregation of the pattern at distance r , and $O(r) < \lambda$ regularity.

Edge correction

Edge effects may arise in calculating point-pattern statistics due to the fact that data points lying outside the study region (ones that could potentially influence the pattern inside the study region) have not been

sampled and are unknown. This means that sample circles or rings used in calculating point-pattern statistics may fall partially outside the study region and will produce a biased estimate of the point-pattern unless a correction is applied. One method used to avoid edge effects is to sample an additional buffer zone, with width r equal to the largest scale used in the analysis, surrounding the main study area. Only points lying inside the main study area are utilized as centers in calculating the point-pattern statistics (Haase 1995). Clearly, the shortcoming of this method is that only the points within the inner plot can be analyzed and for large scales a large buffer zone must be utilized. For rectangular study regions, a second method using a toroidal edge correction can be employed to avoid edge effects. This involves replicating the observed pattern eight times and then surrounding the original pattern with the eight copies to form a 3×3 array (Ripley 1979, 1981, Upton and Fingleton 1985, Haase 1995). The justification of toroidal edge correction is that the observed rectangle represents a random sample of all the rectangles that may have been observed, and therefore the best guess on the appearance of the adjacent rectangles is that they look identical to the sampled rectangle (Upton and Fingleton 1985). Utilizing a buffer zone or toroidal edge correction is only necessary if the degree of edge in the analysis is high and a large proportion of the area sampled around focal points lies outside the main study area, e.g. for transect data or small plots (Haase 1995). However, if most of the area sampled around focal points falls within the study area, a third form of edge correction, employing a weighting that corrects for the proportion of the sample area lying outside the study area, can be utilized as explained below (Ripley 1981, Bailey and Gatrell 1995, Haase 1995, Goreaud and Pélissier 1999).

Analytical and numerical implementation

There are basically two approaches to estimate $K_{12}(r)$ and $O_{12}(r)$ from the data: an analytical approach, and a numerical approach. Analytical approaches use geometric formulas to calculate weights that correct for the area of the circles lying outside the study region (Haase 1995, Goreaud and Pélissier 1999), whereas numeric approaches use an underlying grid of cells for implementation of Eq. 1 and 3 and do not require edge correction.

Analytical approach.

The common analytical estimator for $K_{12}(r)$ was proposed by Ripley (1976, 1981). It is based on all distances d_{ij} between the i th point of pattern 1 and the j th point of pattern 2 and is given by:

$$\hat{K}_{12}(r) = \frac{A}{n_1 n_2} \sum_{i=1}^{n_1} \sum_{j=1}^{n_2} \frac{I_r(d_{ij})}{w_{ij}} \quad (5)$$

where n_1 and n_2 are the number of points of pattern 1 and 2, respectively, A is the area of the study region, I_r is a counter variable [$I_r(d_{ij}) = 1$ if $d_{ij} \leq r$, and $I_r(d_{ij}) = 0$ otherwise], and w_{ij} is a weighting factor to correct for edge effects. The weight w_{ij} is the proportion of the area of a circle centered at the i th point of pattern 1 with radius d_{ij} that lies within the study region. For example, if only half of the circle falls in the study region, Eq. 5 counts effectively two points for each point encountered in the incomplete circle. This edge correction is based on the assumption that the region surrounding the study region has a point density and distribution pattern similar to the nearby areas within the boundary (Getis and Franklin 1987, Haase 1995) and is in some ways analogous to a toroidal edge correction. Because a point close to the border of the study region is weighted more than a point far away from the border, $\hat{K}_{12}(r)$ may be biased for larger r if narrow and long study regions are analyzed (e.g. transects). The common rule of thumb to avoid this effect is that one should not go to a lag distance longer than half the narrowest dimension. The analytical estimator of the mark-correlation function $g_{12}(r)$ is the analogue to Eq. 5, but the counter variable must account for a ring with width w : $I_r^w(d_{ij}) = 1$ if $r - w/2 \leq d_{ij} \leq r + w/2$ and $I_r(d_{ij}) = 0$ otherwise.

Precise formulas of w_{ij} depend on the shape of the study region and on the location of point i in relation to the boundaries. The derivation of analytical formulas for w_{ij} sometimes requires quite complex algorithms and can be computationally intensive. Ripley (1982), Haase (1995), and more recently Goreaud and Pélissier (1999) reviewed current formulas for edge correction of Ripley's K . Because of the complexity of analytical formulas of w_{ij} , mostly simple circular or rectangular study regions have been used for experimental plots. However, shapes that are more complex are sometimes necessary because of obstacles in the study site, or it may be necessary to omit some parts of a heterogeneous study region to obtain a homogeneous pattern. The shape of the final study region can thus be very complex. Only recently, Goreaud and Pélissier (1999) proposed a general method to deal analytically with study regions of complex shape, by excluding triangular surfaces from rectangular or circular initial shapes.

Numerical approach

Numerical methods require division of the study region into a grid of cells (Fig. 1). Selection of an appropriate cell size is constrained by the sampling error of the coordinates of the points that defines a minimum cell size, and by computational time for larger grids. A resolution coarser than the sampling error can be selected; this will depend on the minimum resolution

of distance classes necessary for responding to the scientific question.

A numerical estimator of $K_{12}(r)$ could determine the weights of Eq. 5 by using the underlying grid. However, the numerical method allows for a slightly different approach that does not "look" outside the study region and therefore does not require edge correction.

This can be achieved by dividing the mean number of points within circles by the mean area of these circles, but counting only points and area inside the study region:

$$\lambda_2 \hat{K}_{12}(r) = \pi r^2 \frac{\frac{1}{n_1} \sum_{i=1}^{n_1} \text{Points}_2[C_i(r)]}{\frac{1}{n_1} \sum_{i=1}^{n_1} \text{Area}[C_i(r)]} \quad (6)$$

$C_i(r)$ is the circle with radius r centered on the i th point of pattern 1, n_1 the total number of points of pattern 1 in the study region, the operator $\text{Points}_2[X]$ counts the points of pattern 2 in a region X , and the operator $\text{Area}[X]$ determines the area of the region X . To implement Eq. 6 we marked each cell (x, y) with an identifier $S(x, y)$ [$S(x, y) = 1$ if the cell with coordinates (x, y) is inside the boundaries of the study region, otherwise $S(x, y) = 0$] and with two additional marks $P_1(x, y)$ and $P_2(x, y)$ that give the number of points of pattern 1 and pattern 2 lying within the cell, respectively. Using these definitions, the numerator of Eq. 6 becomes:

$$\text{Points}_2[C_i(r)] = \sum_{\text{all } x} \sum_{\text{all } y} S(x, y) P_2(x, y) I_r(x_i, y_i, x, y) \quad (7)$$

where (x_i, y_i) are the coordinates of the i th point of pattern 1, and the counter variable I_r defines the circle with radius r that is centered at the i th point of pattern 1:

$$I_r(x_i, y_i, x, y) = \begin{cases} 1 & \text{if } \sqrt{(x - x_i)^2 + (y - y_i)^2} \leq r \\ 0 & \text{otherwise} \end{cases} \quad (8)$$

The denominator of Eq. 6 is calculated analogously to Eq. 7, but it counts cells instead of points:

$$\text{Area}[C_i(r)] = z^2 \sum_{\text{all } x} \sum_{\text{all } y} S(x, y) I_r(x_i, y_i, x, y) \quad (9)$$

where z^2 is the area of one cell. Because Eq. 7 and 9 include the identifier $S(x, y)$ of the study region, only points and cells are counted that are inside the boundaries of the study region. Therefore, the study region can be of any complex shape accommodated by the underlying grid. Using Eq. 6, our numerical estimator of the L -function is given by:

$$\hat{L}_{12}(r) = \sqrt{\frac{K_{12}(r)}{\pi}} - r = r \left(\sqrt{\frac{1}{\lambda_2} \left(\frac{\lambda_2 K_{12}(r)}{\pi r^2} \right)} - 1 \right)$$

$$= r \left(\sqrt{\frac{A}{n_2} \left(\frac{\sum_{i=1}^{n_1} \text{Points}_2[C_i(r)]}{\sum_{i=1}^{n_1} \text{Area}[C_i(r)]} \right)} - 1 \right) \quad (10)$$

where A is the area of the study region, and n_2 the number of points of pattern 2 inside the study region.

The analogous numerical estimate for the bivariate O-ring statistic is:

$$\hat{O}_{12}^w(r) = \frac{\sum_{i=1}^{n_1} \text{Points}_2[R_i^w(r)]}{\sum_{i=1}^{n_1} \text{Area}[R_i^w(r)]} \quad (11)$$

where $R_i^w(r)$ is the ring with radius r and width w centered in the i th point of pattern 1. The numerator and the denominator and of Eq. 11 are the same as given in Eq. 7 and 9, respectively, but the counter variable I_r for circles has to be replaced by a counter variable I_r^w that defines a ring with radius r and width w around the i th point with coordinates (x_i, y_i) :

$$I_r^w(x_i, y_i, x, y) = \begin{cases} 1 & \text{if } r - \frac{w}{2} \leq \sqrt{(x - x_i)^2 + (y - y_i)^2} \leq r + \frac{w}{2} \\ 0 & \text{otherwise} \end{cases} \quad (12)$$

Methods for delineating homogeneous sub-regions

Analysis of heterogeneous point patterns is difficult because most methods related to Ripley's K-function have been developed for homogeneous point patterns. The classical exploratory approach for univariate point patterns is to compare a given point pattern to the null model of a homogeneous Poisson process which generates patterns consistent with complete spatial randomness (CSR). However, for heterogeneous patterns this null model is not appropriate because first-order effects may interfere with second-order effects. In the examples section ("Virtual aggregation and bias in the univariate L-function") we show e.g. that analysis of univariate point patterns with larger gaps (i.e. areas without points or with a low density of points) can lead to severe misinterpretation of the spatial structure of the pattern if the null model is CSR. In the case of heterogeneous patterns, CSR might be represented by a heterogeneous Poisson process, a Cox process, or a

Poisson cluster process (Diggle 1983, Upton and Fingleton 1985, Bailey and Gatrell 1995). However, the corresponding mathematical tools for analytical implementation of these null models are complicated (Batista and Maguire 1998, Pélissier and Goreaud 2001). One possibility to retain the methods for homogeneous patterns for the analysis of heterogeneous patterns is to define homogeneous sub-regions and to analyze the spatial structure within these separately (Pélissier and Goreaud 2001). This requires methods for delineating irregularly shaped sub-regions within heterogeneous patterns, or methods for detection of clusters or gaps in point patterns.

Dale and Powell (2001) presented an approach for detecting gaps and clusters in point patterns that is closely related to Ripley's K-function. They did not center the circles on points of the pattern (as is done when calculating the K-function), but on circles that go through all three points of any trio of points in the map (circumcircle method). Circles that include large gaps will have many fewer observed than expected points and the circles that neatly enclose patches will have many more than expected (Dale and Powell 2001). Another method to detect and measure clusters in spatially referenced count data was presented in Perry et al. (1999). This method works through equating the degree of spatial pattern in an observed arrangement of counts to the minimum effort that the individuals in the population would need to expend to move to a completely regular arrangement in which abundance was equal in each sample unit. In practice, this effort is equated with the minimum distance required to move to complete regularity (Perry et al. 1999).

A simple exploratory approach that tests for homogeneity relies on the fact that the number of points in plots of size W of homogeneous point patterns follows a Poisson distribution. For a point pattern with intensity λ a test for homogeneity involves comparison of the estimated frequency distribution $\hat{P}_k(\lambda W)$ [of finding k points in an arbitrary plot of size W] with the expected frequency distribution under homogeneity, a Poisson distribution

$$P_k(\lambda W) = \frac{(\lambda W)^k}{k!} e^{-\lambda W} \quad (13)$$

with mean λW . Depending on the kind of heterogeneity, there are different possibilities to use a combined analysis of expected and estimated frequency distribution together with an estimate of the first-order intensity for delineation of homogeneous sub-regions. We discuss two cases of simple heterogeneity.

Detecting a mixture of two Poisson processes

If a point pattern comprises two internally homogeneous sub-regions with different intensities, the estimated

frequency distribution will show two overlapping peaks (Fig. 2A). Pélissier and Goreaud (2001) proposed visual inspection of $\hat{P}_k(\lambda W)$ to find the critical value of k ($=k_{sep}$) that separates both Poisson distributions. In a next step, they used spatial interpolation techniques to approximate the heterogeneous first-order intensity λ and constructed contour lines with the critical intensity $\lambda_{sep} = k_{sep}/W$ to delineate the two homogeneous sub-regions.

Detecting gaps

If the point pattern is internally homogeneous but contains a gap (i.e. a larger region with low density of points), a plot of size W may be part of the gap if it contains less points than one would expect under homogeneity. The minimal number of points k_{min} in the plot that would be still probable under homogeneity can be estimated by accumulating the “left tail” of the expected Poisson distribution until a small value p_0 is reached:

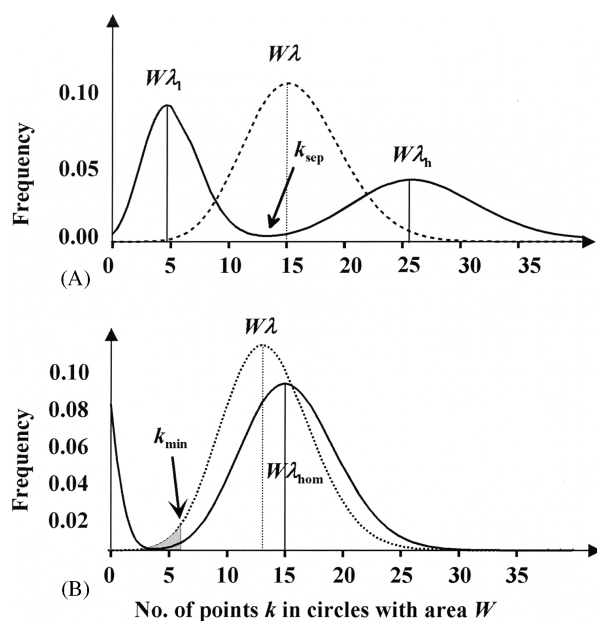


Fig. 2. Delineating homogeneous sub-regions for cases of simple heterogeneity. (A) Two sub-regions with different first-order intensities λ_l and λ_h ($\lambda_l = 0.01$, $\lambda_h = 0.05$) of equal size. The broken line shows the expected frequency distribution of the number of points in plots of size W under overall homogeneity, and the solid line the frequency distribution resulting from the mixture of the two Poisson processes. The arrow indicates the value of k_{sep} that separates the two sub-regions (Pélissier and Goreaud 2001). (B) Homogeneous pattern with a gap. The solid line shows the frequency distribution that results from a homogeneous pattern with $\lambda_{hom} = 0.03$ that contains a gap covering 12.5% of the study region. The broken line shows the expected frequency distribution under overall homogeneity ($\lambda = 0.026$). Under a homogeneous pattern with $\lambda = 0.026$, circles of size $W = 500$ would contain 5 or less points with probability 0.009, thus $k_{min} = 6$ for $p_0 = 0.01$.

$$\sum_{k=0}^{k_{min}-1} P_k(\lambda W) \leq p_0 \quad \text{and} \quad \sum_{k=0}^{k_{min}} P_k(\lambda W) > p_0 \quad (14)$$

Thus, there is only a probability p_0 that a plot of size W contains fewer than k_{min} points, and a plot that contains $< k_{min}$ points is probably part of a gap. An estimate of the first-order intensity $\lambda(x, y)$, and contour lines with $\lambda_{min} = k_{min}/W$ can then be used to separate the gaps from the homogeneous region. Theoretically, this approach does not require an estimate of $\hat{P}_k(\lambda W)$, but in practice, $\hat{P}_k(\lambda W)$ is needed as an exploratory tool (see examples section “Delineating homogeneous sub-regions”). Note that this approach can be applied analogously for detecting clusters. In this case, one has to consider the upper tail of the Poisson distribution instead of the lower tail.

Implementation

Formal statistical tests of the estimated frequency distribution $\hat{P}_k(\lambda W)$ against its theoretical distribution require independence of the sample (i.e. non-overlapping plots) for estimation of $\hat{P}_k(\lambda W)$. Because we aim to remove larger-scale variation in the intensity $\lambda(x, y)$ we will tend to select larger plot sizes W . On the other hand, larger non-overlapping plots will produce smaller sample sizes, which can be a problem if the study region is small. A pragmatic, but less rigorous, approach for exploratory analysis is the use of overlapping plots. However, this requires some methods to account for non-independence. For example, Pélissier and Goreaud (2001) proposed to use a buffer zone between the two homogeneous sub-regions to account for unclear transition between the dense and sparse parts of the study region.

The numerical approach suggests a simple estimator of the non-constant first-order intensity $\lambda(x, y)$:

$$\hat{\lambda}^R(x, y) = \frac{\text{Points}[C_{(x,y)}(R)]}{\text{Area}[C_{(x,y)}(R)]} \quad (15)$$

where $C_{(x,y)}(R)$ is a circular moving window with radius R that is centered in cell (x, y) . This is basically a kernel estimate with fixed bandwidth R (Diggle 1985, Bailey and Gatrell 1995). As edge correction, the number of points in an incomplete circle is divided by the proportion of the area of the circle that lies within the study region. Using Eq. 15 has the advantage that a fitting procedure for contour lines is not required, and the moving window procedure can easily be performed for several spatial scales R .

A rigorous method for detecting gaps would place e.g. adjacent (rectangular) plots over the study region and remove all plots containing fewer points than k_{min} .

However, the plot size W must be sufficiently large, otherwise gaps cannot be distinguished from empty plots that would occur under homogeneity with a probability $P_0(\lambda W)$. Equation (14) can be used to calculate the plot size $W_{\min}$ for which empty plots will occur with (a low) probability p_0 :

$$P_{k=0}(\lambda W) = e^{-\lambda W_{\min}} = p_0 \text{ which yields } W_{\min} = -\ln(p_0)/\lambda \quad (16)$$

Non-overlapping plots of area $W > W_{\min}$ will make the delineation of gaps very coarse if the overall density $\lambda = n/A$ is low (n is the number of points in the study region with area A). As a pragmatic alternative, we propose in the examples section “Delineating homogeneous sub-regions” a combined analysis of the estimated frequency distribution $\hat{P}_k$ and the moving-window estimate $\hat{\lambda}^R$ of $\lambda(x, y)$.

Null models

The keys for successful application of Ripley’s K -function and the O -ring statistic are the selection of an appropriate null model that responds to the specific biological question asked, and correct interpretation of a given departure of the data from the null model. One approach to find an appropriate null model is based on the mathematical form of $K(r)$ and $g(r)$, which are known explicitly (or as an integral) for a number of potentially useful classes of spatial point processes (Ripley 1981, Diggle 1983, Upton and Fingleton 1985, Bailey and Gatrell 1995, Dixon 2002). This is a two-stage process. First, inspection of the estimated $\hat{K}(r)$ may suggest plausible models for the underlying point process, and the parameters that control the process can be fitted through comparison of the expected and estimated K -function (e.g. Diggle 1983, Batista and Maguire 1998, Dixon 2002). Second, approximate confidence envelopes for the K -function based on the fitted models are constructed by Monte Carlo simulations of the stochastic process with the parameter estimates obtained by the fitting procedure. This approach has mainly been adopted by statisticians interested in spatial point processes, but the mathematical tools can be complicated, especially for parameter fitting and correcting edge effects. As a consequence, ecologists have mostly used the simplest case only, the null model of CSR (but see Diggle 1983, Batista and Maguire 1998). The numerical approach facilitates simple implementation of a variety of null models that e.g. account for first-order heterogeneity, or are adapted to specific biological questions. Because there are fundamental differences between the univariate and the bivariate case, we will review null models separately for the univariate and the bivariate case.

Null models for univariate point pattern

Complete spatial randomness

The simplest and most widely used null model for univariate point patterns is complete spatial randomness (CSR) that can be implemented as a homogeneous Poisson process. Homogeneous means that the first-order intensity λ is constant over the study region (there are no first-order effects), and Poisson means that the probability of finding k points in an area W follows a Poisson distribution with mean λW . Thus, any point of the pattern has an equal probability of occurring at any position in the study region, and the position of a point is independent of the position of any other point (i.e. points do not interact with each other). Due to practical problems with edge correction, CSR has mostly been applied in study regions of simple rectangular or circular shape (but see Goreaud and Pélissier 1999). If a homogeneous pattern is spatially restricted by obstacles or environmental heterogeneity (e.g. differences in soil), the appropriate null model is CSR, but applied only within an irregularly shaped study region. In the examples section “Virtual aggregation and bias in the univariate L -function” we show that in this case application of CSR in a rectangular study region that encompasses the pattern can lead to severe misinterpretation of the second-order structure of the pattern. If appropriate software for edge correction of irregularly shaped study regions is not available, smaller regularly shaped sub-regions of the pattern have to be analyzed. However, this may reduce the sample size considerably. Note that the numerical approach (Eq. 7 and 9) can deal with any irregularly shaped study region accommodated by the underlying grid.

Heterogeneous Poisson process

If a pattern is not homogeneous, the null model of CSR is not suitable for exploration of second-order characteristics, and a null model accounting for first-order effects has to be used to reveal “true” second-order effects. The heterogeneous Poisson process is the simplest alternative to CSR if the pattern shows first-order effects. The constant intensity of the homogeneous Poisson process is replaced by a function $\lambda(x, y)$ that varies with location (x, y) , but the occurrence of any point remains independent of that of any other. The intensity function $\lambda(x, y)$ determines the process completely, and numerical implementation of this null model is a matter of finding an appropriate estimate of the intensity function.

The numerical approach suggests a simple method to implement the heterogeneous Poisson process using the moving-window estimate $\hat{\lambda}^R$ of the intensity function $\lambda(x, y)$ (Eq. 15): a provisional point is placed at a random cell (x, y) in the study area, but this point is only retained with a probability given through $\hat{\lambda}^R(x, y)$. This

procedure is repeated until n points are distributed. The moving window estimator $\hat{\lambda}^R(x, y)$ involves a decision on an appropriate radius R of the moving window. Because the bandwidth R is the scale of smoothing, possible departure from this null model may only occur for scales $r < R$, and for small moving windows it will closely mimic the original pattern, whereas a large moving window approximates CSR.

Note that there will always be a subjective component involved in the decision as to whether or not, and at what scales, the pattern is heterogeneous. In general, this decision depends on spatial scale: as compared with the size of the study region, fine-scale variations are generally considered as elements of structure and broad-scale variations as heterogeneity (Pélissier and Goreaud 2001). In some cases, the nature of the data and the strength of trends in the observed pattern may make such judgment relatively straightforward. In other cases, this may be difficult and open to debate and interpretation.

Typical biological situations for application of a heterogeneous Poisson process are presence of exogenous factors (e.g. soil, topography, rocks, etc.) or obstacles that cause irregularly shaped study regions. In fact, a simple variant of the heterogeneous Poisson process can be used as alternative to avoid edge correction for homogeneous point patterns in irregularly shaped study regions: the intensity $\lambda(x, y)$ is zero outside the study region and constant inside.

Random labeling

Random labeling is a somewhat different approach to correct for underlying environmental heterogeneity that can be used where a “control” pattern is available to act as surrogate for the varying environmental factor. The assumption of univariate random labeling is that the pattern of controls was created by the same stochastic process as the primary pattern (“cases”). Therefore, the n_1 cases represent a random sub-sample of the joined pattern of the n_2 control points and n_1 case points. The test is devised by computing the univariate K-function for the observed cases, then randomly re-sampling sets of n_1 points from the $(n_1 + n_2)$ points of the cases and controls to generate the confidence limits. Note that this null model makes sense only if there are many more controls than cases. Univariate random labeling is closely related to bivariate random labeling (see below) and has been applied to investigate competitive thinning (Kenkel 1988, Moeur 1993, Batista and Maguire 1998) under the null hypothesis that the survivors are no different from a random draw of the initial cohort.

Poisson cluster process

The Poisson cluster process explicitly incorporates a clustering mechanism. Parent events form a CSR process

and each parent produces a random number of offspring according to a probability distribution $f(\cdot)$. Offspring are spatially distributed around their parent according to some bivariate probability density $g(\cdot)$. The final pattern consists of the offspring only. To avoid edge effects, the parents must be simulated over a region larger than the study region but the offspring falling outside the study region are lost (Bailey and Gatrell 1995). If the number of offspring follows a Poisson distribution and the location of the offspring, relative to the parent individual, have a bivariate, Gaussian distribution, the offspring follow a Neyman-Scott process (Diggle 1983, Cressie 1991, Batista and Maguire 1998, Dixon 2002). The K-function and the pair-correlation function for the Neyman-Scott process are given by:

$$K(r, \sigma, \rho) = \pi r^2 + \frac{1 - \exp(-r^2/4\sigma^2)}{\rho}$$

$$g(r, \sigma, \rho) = 1 + \frac{\exp(-r^2/4\sigma^2)}{4\pi\sigma^2\rho} \quad (17)$$

where ρ is the intensity of the parent process, and σ^2 the variance of the Gaussian distribution. Because σ is the standard deviation of the distance between each offspring and its parents, the cluster size yields $\sim 2\sigma$. For scales r below the cluster size (i.e. $r < 2\sigma$) the K-function can be approximated by $K(r) = r^2\pi + r^2/(4\rho\sigma^2)$ (Diggle 1983), and the L-function is approximated by

$$L(r, \sigma, \rho) \approx r \left(\sqrt{1 + \frac{1}{4\pi\rho\sigma^2}} - 1 \right) \quad (18)$$

Because the parameters ρ and σ are unknown, they must be fit by comparing the empirical $\hat{K}(r)$ with the theoretical K-functions (Diggle 1983, Batista and Maguire 1998, Dixon 2002). Rough initial estimates for ρ and σ can be obtained by using two properties of the L-function: the maximal value of $\hat{L}(r)$ occurs for a value of r slightly above the cluster size 2σ , and the L-function increases almost linearly for $r < 2\sigma$ with slope given in Eq. 18.

Hard-core process

A hard-core process is the simplest extension of CSR to describe small-scale regularity. Dixon (2002) reviews analytical formulas of the K-function under different hard-core processes. For numerical simulation of a hard-core process, CSR is “thinned” by deletion of all pairs of points a distance less than δ apart. A “soft-core” variant of this inhibition process distributes provisional points due to CSR and retains a point that is closer than distance δ from an already accepted point with a probability that varies within distance d ($< \delta$) between 1 at $d = \delta$ and 0 at $d = 0$. A typical biological situation for application of a hard-core process is a hypothesized

negative effect of competition that may work only up to a certain distance d . A hard-core process is also appropriate if spatially extended objects, such as non-overlapping plants of finite size, are analyzed. In this case, the usual point-approximation will introduce a bias because the plants are constrained by their diameters to be at least a certain distance apart. As a consequence, the point approximation will indicate e.g. regularity but this appearance of regularity may conceal significant small-scale aggregation (Simberloff 1979, Prentice and Werger 1985).

Null models for bivariate point pattern

Interpreting a bivariate K-function or O-ring statistic can be confusing because it differs from the univariate case. In the univariate case, visualization of the pattern can provide an intuitive idea of the first and second-order properties of a pattern. However, in the bivariate case we analyze the spatial relation between two spatial patterns at different spatial scales where each pattern individually can have a complicated spatial structure. Confusion may also arise because there is not one simple and intuitive null model such as CSR, and because a null model based on CSR (i.e. randomization of both patterns) leads to an inadequate test of the bivariate pattern.

The relationship between two patterns can be contrasted to two conceptually different null models: independence and random labeling. The null model of random labeling assumes that both patterns were created by the same stochastic process, and each of the two patterns taken separately represents random “thinning” of the joined pattern. In contrast, the null model of independence assumes that the two patterns were generated by two independent processes (e.g. one process generated the locations of shrubs, and the other process generated the locations of grass tufts). Departure from independence indicates that the two processes display attraction or repulsion, regardless of the univariate pattern of either group by itself. The distinction between independence and random labelling requires some care and consideration (Dixon 2002). When there is no relationship between two processes, the two approaches lead to different expected values of $K_{12}(r)$ and $O_{12}(r)$, and to different procedures for generating null models. Failure to distinguish between random labelling and independence may lead to the analysis of data by methods which are largely irrelevant to the problem at hand (Diggle 1983). Random labelling and independence are equivalent only if all the component processes are homogeneous Poisson processes.

Independence

Testing for independence is more difficult than testing for CSR in the univariate case because inferences are conditional on the second-order structure of each pattern (Dixon 2002). This is because the theoretical values of $K_{12}(r)$ and $O_{12}(r)$ do not depend on CSR of the component patterns and therefore no assumption can be made about models for either of the component patterns. Thus, the null model of CSR is not appropriate to test for independence; the separate second-order structures of the patterns need to be preserved in their observed form in any simulation of the null model, but one has to break the dependence between the two patterns. One way of achieving this is by simulations that involve random shifts of the whole of one component pattern relative to the other. In practice, a rectangular study region is treated as a torus where the upper and lower edges are connected and the right and left edges are connected.

Random labeling

In the case of random labelling we ask not about the interaction between two processes, but about the process that assigns labels to points, conditioning the observed locations of the points of the joined pattern. Therefore, the null model of random labelling requires no assumptions about the specific form of the two underlying component processes. Because both component patterns taken separately represent “random thinning” of the joined pattern a numerical implementation of random labelling involves repeated simulations using the fixed $n_1 + n_2$ locations of pattern 1 and 2, respectively, but randomly assigning “case” labels to n_1 of these locations (Bailey and Gatrell 1995). From their definition, K-functions (and g-functions) are invariant under random thinning and therefore we would expect $K_{12}(r) = K_{21}(r) = K_{11}(r) = K_{22}(r)$. This suggests that a useful way of investigating departures from random labelling is to assess the significance of differences amongst estimates of these functions (Bailey and Gatrell 1995). Each pairwise difference evaluates different biological effects. The difference $\hat{K}_{11}(r) - \hat{K}_{12}(r)$ for example evaluates whether points of type 1 tend to be surrounded by other points of type 1, while $\hat{K}_{11}(r) - \hat{K}_{22}(r)$ evaluates whether one pattern is more (or less) clustered than the other (Dixon 2002).

Typical biological situations for the application of bivariate, random labelling are cases where an underlying pattern is imposed on both patterns, i.e. a heterogeneous environment. Random labelling has not been applied much in ecological applications (but see Dixon 2002), but it is frequently used in the epidemiological context to account for the natural variation in the background population

(i.e. “population at risk”, Bailey and Gatrell 1995).

Space-time clustering

Except for random thinning, all null models discussed so far are concerned with the properties of point patterns without explicit reference to time. However, we might be interested in how spatial patterns change over time – that is, whether events cluster in space over time. For example, in plant ecology the locations of plants in a study region might be re-sampled with a certain time lag e.g. for determining mortality and recruitment rates or for studying successional patterns. In order to assess the spatio-temporal relationships among plants, one can regard the vegetation maps at different times as being different patterns and analyze these spatio-temporal patterns using the bivariate K-function or O-ring statistic. Space-time clustering has been investigated e.g. by Wiegand et al. (1998) in an analysis of a model that simulated vegetation dynamics of colonizer and successor species in a South African shrubland. Patterns of space-time clustering are of particular interest in this system where cyclic succession produced strong time lags in the establishment of successor species dependent upon the earlier establishment by colonizer species. Especially in an epidemiological context, spatial data on e.g. occurrence of a disease might not stem from snap-shots of the disease at different times, but occurrence might be continuously mapped with a time label attached. In this case, one can define the ‘space-time’ K-function $K(r, t)$ analogously to Eq. 1. If the processes operating in time and space are independent, $K(r, t)$ should be the product of separate space and time K-functions (Bailey and Gatrell 1995).

Antecedent conditions

In some cases antecedent conditions may influence the choice of an appropriate null model. For example, in space-time clustering we need to keep the locations of the earlier pattern fixed, and randomize only the later pattern following an appropriate null model. Similarly, for investigating the relationship between adult trees (pattern 1) and seedlings (pattern 2) an appropriate null model to test for repulsion or attraction would be to randomize the locations of the seedlings (because they could potentially be found at the entire study region) and to keep the locations of the trees fixed. Randomizing the locations of the trees would be inappropriate because they did not change their position during the development of the seedlings. Moreover, possible repulsion or attraction between seedlings and trees might be obscured by randomizing the locations of the trees.

Examples

Virtual aggregation and bias in the univariate L-function

If a pattern is not homogeneous, the null model of CSR is not suitable for exploration of second-order characteristics. This is because large-scale, first-order effects introduce a systematic bias in the univariate K-function, not only at larger scales, but also at smaller scales. In this case, an observed departure from CSR could well be due to first order effects rather than to second order effects (Bailey and Gatrell 1995). This can be understood intuitively, when imagining a point pattern that comprises a single internally homogeneous cluster in the center of the study region (Fig. 3A). In this case the local density of points in the cluster will be higher than the overall density of points in the entire study region. As a consequence, there are always more points in the closer neighborhood of other points than expected under homogeneity, and the K-function will indicate aggregation at smaller scales even if the pattern is random inside the cluster. We call this phenomenon “virtual aggregation.”

To demonstrate this intuitive idea mathematically, we imagine a univariate point pattern with overall intensity λ that forms an internally random cluster covering the proportion c of the study region. There are no points outside the cluster. Because sub-regions of the cluster satisfy CSR, the probability $O(r)$ of finding a point at the closer neighborhood r of other points will be approximately constant, i.e. $O(r) = g\lambda$ with $g = 1/c$. To obtain the corresponding K-function we integrate Eq. 4 using $g(r) = O(r)/\lambda = g$ (Eq. 3) and obtain $K(r) = \pi g r^2$, which yields:

$$L(r) = r(\sqrt{g} - 1) \quad (19)$$

Thus, under virtual aggregation we observe an L-function that increases at smaller scales linearly, and the extent of virtual aggregation, given through the slope $\sqrt{g} - 1$, is inversely related to the fraction c of the study region covered by the cluster. Note that for smaller scales (i.e. scales r below the cluster size) the functional form of $L(r)$ under virtual aggregation is the same as under a Neyman-Scott process (Eq. 18 and 19). This is not surprising because virtual aggregation is caused by clustering. The difference is that the cluster size under virtual aggregation is defined to be large, while the Neyman-Scott process can be applied for any cluster size.

The L-function can increase under virtual aggregation only over a limited range of scales; it will start to drop if a notable proportion of circles overlap the part of the study region outside the cluster. Finally, the L-function will approach zero for very large scales r because then all

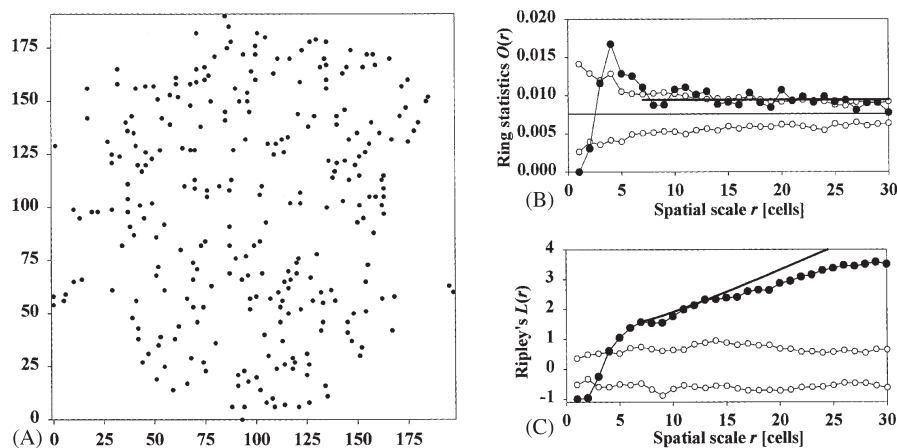


Fig. 3. Univariate analysis of a point pattern, and comparison between Ripley's L-function and the O-ring statistic. (A) map showing the rectangular study region and the point pattern. The study region was divided into a grid of 198×191 cells, and the pattern comprises $n = 287$ points. The overall intensity λ of the pattern in the rectangular study region is $\lambda = 0.0076$. (B) the O-ring statistic (solid circles), giving the local neighborhood density of the pattern at scale r and confidence envelopes (open circles). Confidence envelopes are the highest and lowest $O(r)$ of 99 randomizations of the pattern over the study region. The solid line gives the intensity λ of the pattern in the rectangular study region. The bold line is the average of the O-ring function for scales $r = 7 - 30$. The ring width was one cell unit. (C) Ripley's L-function with confidence envelopes (open circles) constructed in the same way as in (B). The bold line is the approximation Eq. 21 of the L-function for $g = 1.34$ at scales $r = 7 - 30$.

points will be located within each circle, i.e. $K(r) = \pi r^2$, and $L(r) = 0$.

If the pattern shows virtual aggregation but additionally true second-order effects (i.e. a non constant pair-correlation function $g(r)$ at scales $r < r_1$, and $g(r) = g$ for $r > r_1$), integration of Eq. 4 yields

$$K(r) = \int_0^{r_1} 2\pi r' g(r') dr' + \int_{r_1}^r 2\pi r' g dr' \\ = K(r_1) - \pi g r_1^2 + \pi g r^2 \quad (20)$$

and the L-function becomes:

$$L(r) = -r + \sqrt{\frac{K(r_1)}{\pi} - g r_1^2 + g r^2} \quad (21)$$

which collapses back to Eq. 19 if there are no second-order effects (i.e. $K(r_1) = \pi g r_1^2$). Note that Eq. 21 approximates the impact of virtual aggregation only for a limited range of scales r , and for large scales the assumption $g(r) = g$ does not hold because in this case the circles will overlap the gap. Figure 4 shows how second-order effects at small scales (i.e. a given $K(r_1) \neq \pi g r_1^2$, Eq. 21) impact the L-function at higher scales.

Weak virtual aggregation increases the local density $O(r)$ at smaller scales r only slightly and it should therefore not seriously affect the outcome of second-order analysis. However, the problem is that the Monte Carlo test for Ripley's K will indicate highly significant aggregation because the K -function is a cumulative measure where aggregation at smaller scales influences the estimate at larger scales (Eq. 21). The Monte Carlo

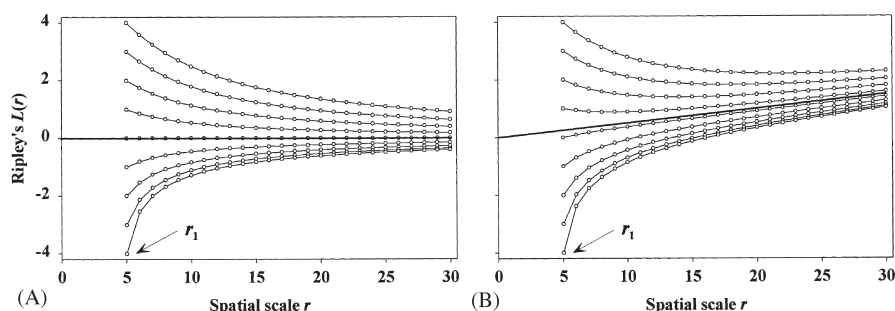


Fig. 4. The memory in the accumulative Ripley's L-function. (A) No virtual aggregation (i.e. $g = 1$ in Eq. 21), but true second-order effects up to scale r_1 , and no second-order effects for scales $r > r_1$. The different curves show Eq. 21 for initial values $L(r_1) = -4, -3, -2, -1, 0, 1, 2, 3$, and 4 . The bold line gives the L-function without second-order effects (i.e. $K(r) = \pi r^2$). (B) Virtual aggregation ($g = 1.2$ in Eq. 21), true second-order effects up to scale r_1 , and no second-order effects for scales $r > r_1$. The bold line gives the L-function without second-order effects (i.e. $K(r_1) = \pi g r_1^2$). Initial values of $L(r_1)$ same as in (A).

test for the non-accumulative O-ring statistic, however, will indicate the expected weak aggregation. Our first example (Fig. 3) illustrates this point. The analysis using the O-ring statistics reveals a non-random pattern at scales $r = 1-6$ due to second-order effects, and virtual aggregation at smaller scales $r = 7-30$ with an approximately constant local neighborhood density $O(r) = g\lambda = 0.0103$ (bold line in Fig. 3B) well above the overall intensity $\lambda = 0.0076$ of points in the entire study region (dashed line in Fig. 3B). As expected, the Monte-Carlo test shows only weak evidence for aggregation. However, Ripley's L indicates strong aggregation at scales $r = 7-30$ with an increasing L-function (Fig. 3C). This result shows that a frequent (implicit) assumption of second-order analysis (univariate Ripley's K will reveal under large-scale heterogeneity the true second-order properties at small scales) is wrong. As we have shown, undetected virtual aggregation leads to an erroneous interpretation of the univariate L-function at smaller scales.

Delineating homogeneous sub-regions

Evidently, the pattern shown in Fig. 3A contains gaps at the edges of the rectangular study region that cause virtual aggregation (Fig. 3B). Now we apply the method for detecting gaps to this pattern. Because our method involves estimation of an empirical frequency distribution $\hat{P}_k$ that gives the number of points in overlapping moving windows, we need to account for the non-independence of the sample. We achieve this by means of a two-step procedure. In a first step we select a large moving window to detect and remove all cells which are clearly gaps, and in a second step we delineate the gaps more closely by selecting a small moving window and by repeating the analysis with the small moving window for the remaining area only.

In the first step we select a large moving window ($R = 50$, one fourth of the width of the study region) to capture large-scale variation in the intensity $\lambda(x, y)$, and to remove cells which are clearly gaps we select additionally a small value $p_0 = 0.01$.

Figure 5D shows the empirical frequency distribution $\hat{P}_k$ for moving windows with $R = 50$ and the corresponding theoretical Poisson distribution. Figure 5D indicates that there are many cells that have in their 50-cell neighborhood less points as expected under CSR, and Eq. 14 yields $k_{\min} = 41$. We remove 5253 cells ($\approx 14\%$ of the rectangular area) corresponding to moving windows containing less than $k_{\min} = 41$ points (dark gray area in Fig. 5A), and 6 isolated points. The overall density in the new (irregularly shaped) study region is $\lambda_1 = (287 - 6)/(198 \times 191 - 5253) = 0.0086$. In the second step we delineate the gaps more closely using the minimal size $W_{\min}$ of a moving window that can discriminate a gap

from a empty plot possible under CSR (Eq. 16). For a small $p_0 = 0.0025$ and $\lambda = 0.0086$ we obtain $W_{\min} = 694$ cells which corresponds to a radius $R_{\min} = 15$. Figure 5E shows the resulting empirical frequency distribution $\hat{P}_k$ for moving windows with $R = 15$. Visual inspection of Fig. 5E suggests removing cells whose moving window contains none or one point. Those are 4368 cells (light gray area in Fig. 5A) and 4 isolated points. Figure 5F shows the empirical and frequency distribution $\hat{P}_k$ for the remaining points in the new study region that now approximates the theoretical frequency distribution reasonably well. The overall density in new study region is $\lambda_2 = (287 - 10)/(198 \times 191 - 9621) = 0.0098$, and in total we remove 25.4% of the initial rectangular study region (thus, $c = 0.746$).

To test the approximation of the O-ring statistics ($O(r) = g\lambda$) and the L-function (Eq. 21) under virtual aggregation we compare the predicted and the observed L and O-functions. Our prediction under virtual aggregation is $O(r) = g\lambda = 0.0102$ ($g = 1/c = 1.34$, $\lambda = 0.0076$). The mean value of $\hat{O}(r)$, taken for scales $r = 7-30$, yields 0.0103 which is in excellent accordance with our prediction (Fig. 3B). The predicted L-function Eq. 21 under virtual aggregation and second-order effects ($\hat{L}(r = 7) = 1.58$) is in good accordance with the observed $\hat{L}$ (Fig. 3C), however, at scales $r > 15$ the observed $\hat{L}$ is lower than the predicted. Clearly, this is because for larger scales a notable proportion of the circles used for calculation of the K-function overlap the gap, and consequently $\hat{L}(r)$ drops.

After detecting gaps we repeated the second-order analyses, but performed the analyses only in the homogeneous sub-region (Fig. 5A), thus excluding the gaps and the 10 isolated points. Figure 5B shows that the first-order effects had a relatively weak impact on the shape of the O-ring function (Fig. 3B, 5B). The only important differences are that the weak aggregation at scales $r > 6$ disappeared and that aggregation at $r = 5, 6$ became weak or non-significant. In contrast, removal of first-order effects had a marked impact on the shape of $\hat{L}(r)$ and on the confidence envelopes (Fig. 3C, 5C). The analysis performed in the initial rectangular study region indicated strong aggregation for $r > 5$ (Fig. 3C) while the analysis in the homogeneous sub-region shows that the pattern is random at scales $r > 5$. The results of both second-order statistics are now in accordance, except for the clear differences caused by the cumulatively property of Ripley's L at scales $r = 2$, and 4.

The pattern in Fig. 3A constitutes a large patch with no points in the edges of the rectangular study region. Now we study the effect of the opposite pattern, a random distribution of points within a rectangular study region, but with a large gap in the middle (Fig. 6A). We created this pattern by distributing points at random over a 101×101 cell grid and then removing all points inside a circle with radius 25 cells in the center of the

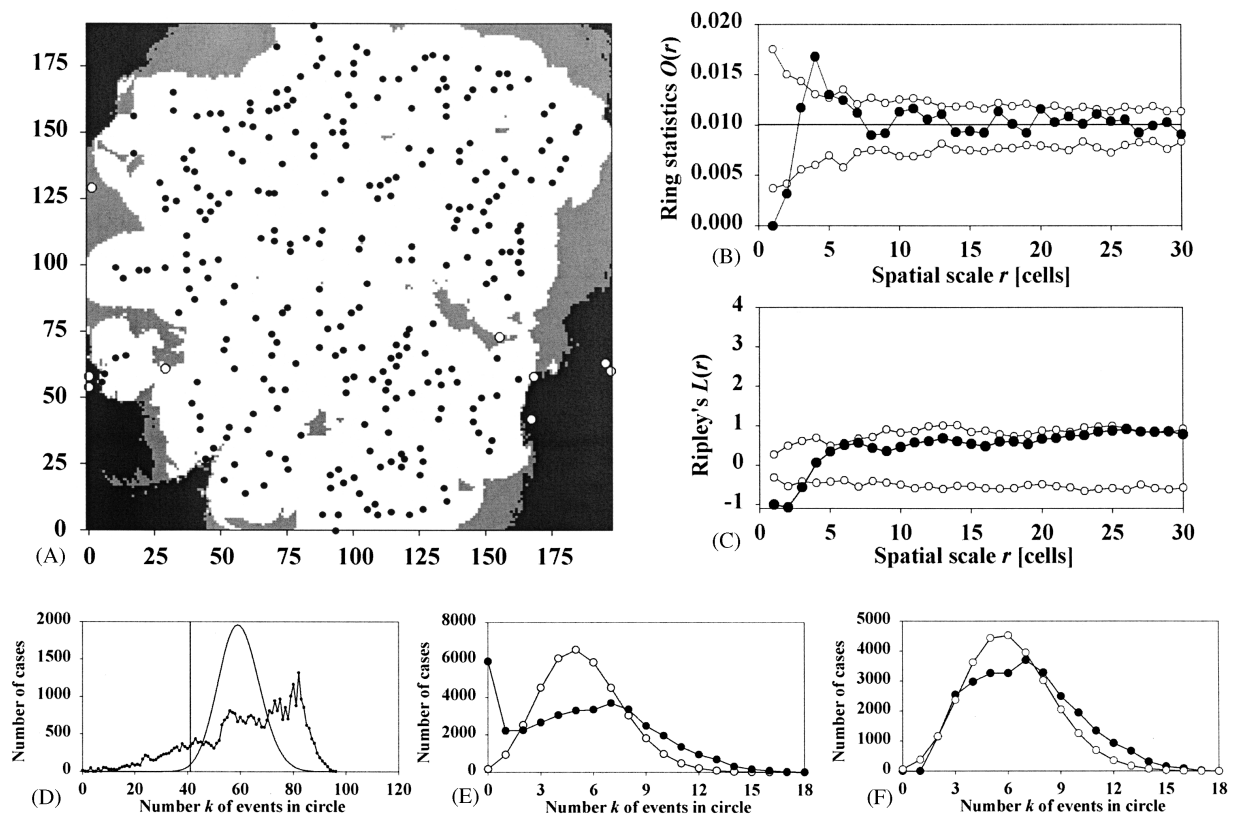


Fig. 5. Delineating the homogeneous sub-region of the study region. (A) Map showing the point pattern and the regions of the study region that were removed during the two-step procedure. Dark grey: cells removed in first step of algorithm to detect gaps ($R = 50$, $p_0 = 0.01$, $k_{\min} = 41$), light gray: cells removed in second step ($R = 15$, $k_{\min} = 2$), white: homogeneous study region. White dots are isolated points that are removed. (B) The O-ring statistic (solid circles) with confidence envelopes (open circles). Confidence envelopes are the highest and lowest $O_{11}(r)$ of 99 randomizations of the remaining 277 points of the pattern over the sub-region where the pattern is homogeneous (white region in Fig. 5A). The solid horizontal gives the mean intensity of the pattern. The ring width was one cell. (C) Ripley's L-function with confidence envelopes (open circles) constructed in the same way as in B. (D) Observed number of points within circular moving windows of radius $R = 50$ (solid circles) taken from the initial rectangular study region, and the expected number under overall homogeneity (solid line). The vertical line indicates $k_{\min} = 41$. (E) Observed number of points within moving windows of radius $R = 15$ (solid circles) after removing the first gaps (i.e. taken from the white and light-grey region in Fig. 5A), and the expected number under homogeneity (open circles). (F) Same as E, but after excluding all gaps (i.e. taken only for the white area in Fig. 5A).

square study region (Fig. 6A). At smaller scales (i.e. $r < 18$) only few rings around the points of the pattern overlap the gap (those in the immediate neighborhood of the gap), thus producing virtual aggregation with an above-mean density of points inside these rings (horizontal line in Fig. 6B). However, as r increases, more and more rings overlap the gap, causing the decline in $\hat{O}(r)$ (Fig. 6B). Therefore, for scales $r > 18$ we observe the opposite phenomenon to virtual aggregation, "virtual repulsion" with below-mean density of points inside the rings. From Eq. 19 it follows that the L-function will first increase due to virtual aggregation, and then drop because of virtual repulsion. This prediction is confirmed by the L-function analysis shown in Fig. 6C. Virtual aggregation produces only marginal significant departures from CSR when tested with the O-ring statistic (Fig. 6B), however, the test with Ripley's K

indicates marginal aggregation for scales $r = 1$, but clearly significant aggregation for scales 8–24 (Fig. 6C).

Next we applied our algorithm to delineate the gap. Application of the same method as described above resulted in moving windows with radius $R = 25$, and $R = 9$ in the first, and the second step, respectively, and visual inspection of Fig. 6G suggests removing cells with empty moving windows. Fig. 6H shows that the empirical and expected frequency distribution $\hat{P}_k$ for the new study regions are in good accordance. The overall density in the initial rectangular study region is $\lambda_2 = 195/(101 \times 101) = 0.019$. The overall density in new study region is $\lambda_2 = (195 - 1)/(101 \times 101 - 1800) = 0.0231$, and in total we removed 18% of the initial rectangular study region (i.e. $c = 0.82$). Our prediction under virtual aggregation is $O(r) = g \lambda = 0.0232$ ($g = 1/c = 1.21$, $\lambda = 0.019$) in the O-ring statistic, and a linear

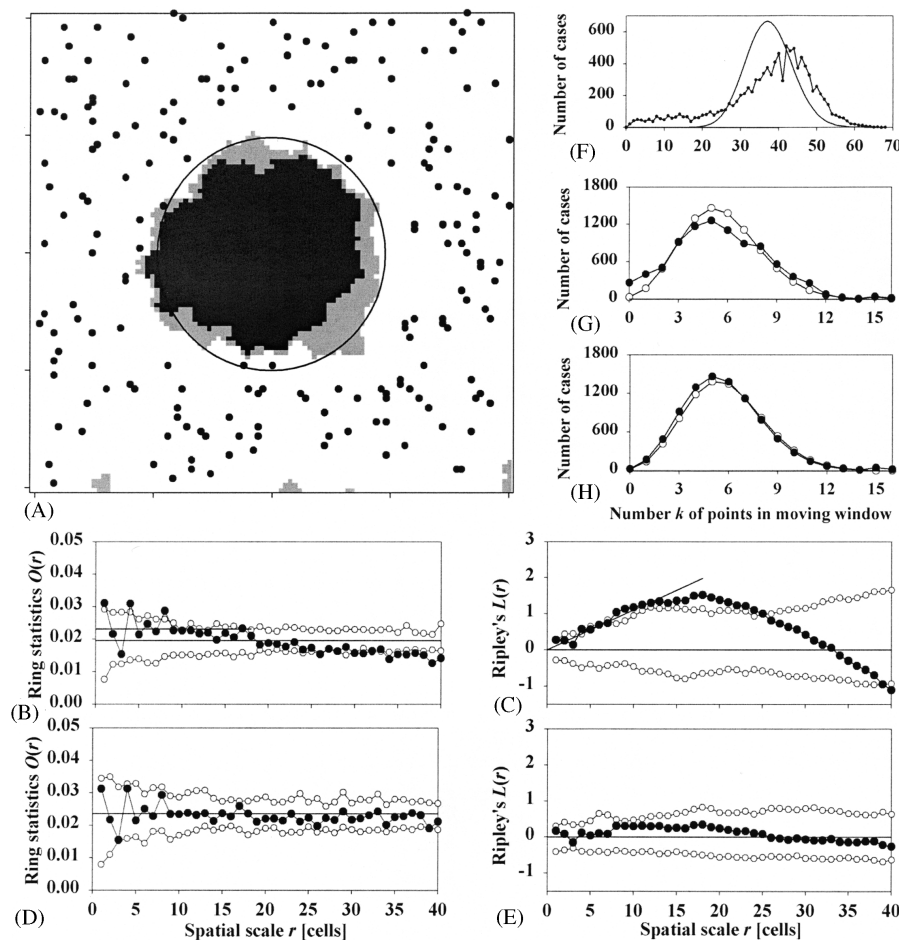


Fig. 6. Analysis of a random pattern with a gap in the center. (A) The pattern was created by randomly distributing 195 points over a 101×101 cell study region, but preventing them to penetrate a circular area in the center (solid circle). Dark grey region: cells removed in the first step of the algorithm for detecting gaps ($R = 25$, $p_0 = 0.01$, $k_{\min} = 23$), light grey: cells removed in the second step ($R = 9$, $k_{\min} = 1$), white: irregularly shaped, but internally homogeneous study region. (B) The O-ring statistic for the pattern (solid circles) with confidence envelopes (open circles) using the highest and lowest $O(r)$ from 99 replicates of a null model where the points of the pattern were randomly distributed over the entire rectangular study region. The horizontal line shows the mean density $\lambda = 0.0191$ of the pattern the original rectangular study region and the bold line indicate virtual aggregation with above-mean density. The ring width was one unit. (C) Ripley's L-function for the pattern (solid circles) with confidence envelopes (open circles) constructed in the same way as in (B). The bold line indicates virtual aggregation with a linearly increasing segment of the L-function. (D) The O-ring statistic (solid circles) with confidence envelopes (open circles) calculated for the irregularly shaped study region (white area in A) and a CSR null model. The solid line gives the mean density $\lambda = 0.0236$ of the pattern in the irregularly shaped study region. (E) Ripley's L-function of the pattern (solid circles) with confidence envelopes (open circles) constructed in the same way as in (D). (F) Observed number of points within moving windows of radius $R = 25$ (solid circles) taken from the initial rectangular study region, and the expected number under overall homogeneity (solid line). (G) Observed number of points within moving windows of radius $R = 9$ (solid circles) after removing the first gaps (i.e. taken from the white and light-grey region in A), and the expected number under homogeneity (open circles). (H) Same as (E), but after excluding all gaps (i.e., taken only for the white area in A).

increase of the L-function with slope $(\sqrt{g} - 1) = 0.11$. The mean value of the estimated $\hat{O}(r)$ (Fig. 6B), taken for scales $r = 1-18$, yields 0.023, and linear regression of $\hat{L}(r)$, taken for scales $r = 1-19$ (Fig. 6C), yields a slope of 0.097. Both values are in excellent accordance with our prediction.

Repeating the second-order analyses in the homogeneous sub-region (Fig. 6A) finally confirmed the known randomness of the pattern and the linear increase of the L-function at smaller scales and the linear decrease at

larger scales caused by the gap disappears (Fig. 6D, E). Also, the observed and the expected frequency distribution of points in circles of radius $R = 10$ (Fig. 6H) were in accordance. Note that our two examples are extreme cases for illustrative purpose and that in cases where a pattern contains several gaps of different sizes the first-order effects of virtual aggregation or repulsion may be obscured because their scales overlap. For analysis of point patterns that show more complicated first-order heterogeneity, we exemplify the use of null models based

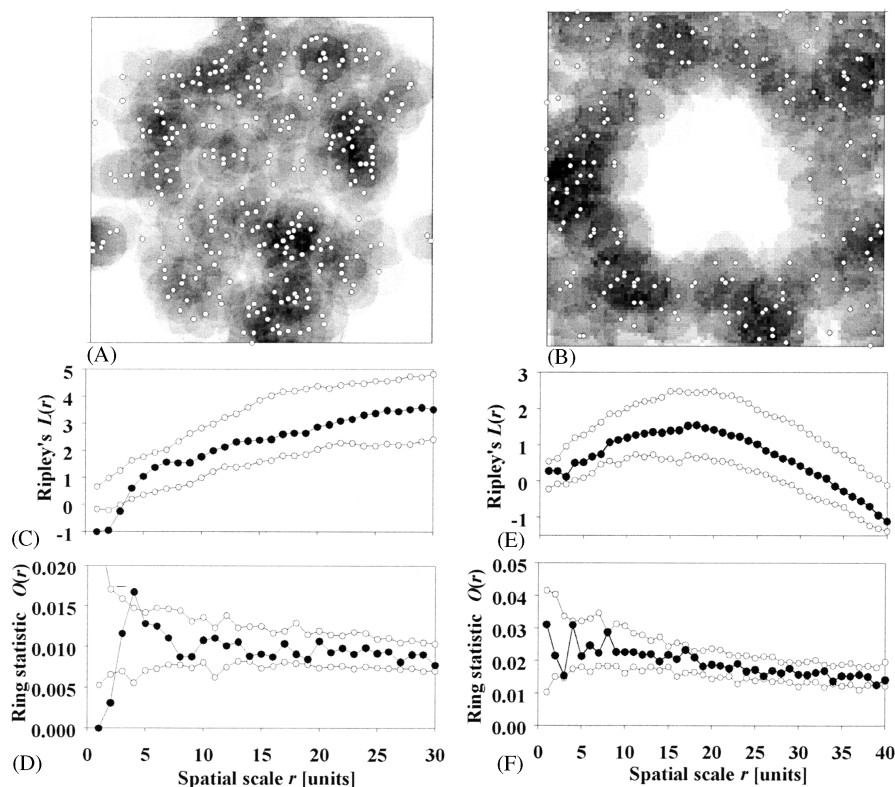


Fig. 7. Second-order statistics and null models based on a heterogeneous Poisson process using the moving window estimate Eq. 15 for approximation of the heterogeneous first-order intensity. (A) Approximation of the first order intensity for the pattern shown in Fig. 3A with a circular moving window of $R = 15$ units. Light gray indicates zero density, and increasing shading indicates increasing density. (B) Approximation of the first order intensity for the pattern shown in Fig. 6A with a circular moving window of $R = 10$ units. (C) Ripley's L-function for the pattern Fig. 3A (solid circles) with confidence envelopes (open circles) using the highest and lowest $O(r)$ from 99 replicates of a null model where the points of the pattern were randomized in accordance to the first-order density shown in (A). (D) O-ring statistic for the pattern Fig. 3A with confidence envelopes constructed in the same way as in (C). The ring width was one unit. (E) Ripley's L-function for the pattern Fig. 6A constructed in the same way as in (C). (F) O-ring statistic for the pattern Fig. 6A constructed in the same way as in (E).

on heterogeneous Poisson processes as tool for analyzing second-order effects despite of the presence of first order heterogeneity.

Null models accounting for full first-order heterogeneity

In this section we illustrate the numerical implementation of the heterogeneous Poisson null model for the two patterns shown in Fig. 3A and 6A. First we calculated the approximate first-order intensity $\hat{\lambda}(x, y)$ of the patterns shown in Fig. 3A (Fig. 7A) and 6A (Fig. 7B) using Eq. 15 and moving windows of $R = 15$ and $R = 10$ cells, respectively. We randomized the points of the patterns in accordance with $\hat{\lambda}(x, y)$. Note that the shape of $\hat{O}(r)$ and $\hat{L}(r)$ does not change under this null model compared to CSR (Fig. 3B, 7D, Fig. 3C, 7C, Fig. 6B, 7F, and Fig. 6C, 7E). How-

ever the confidence envelopes are different because the points of the patterns are randomized in accordance with different null models, and the confidence envelopes for Ripley's L are not symmetric to $L = 0$ (Fig. 7C, E) since the underlying null model differs from CSR.

The results of the second order statistics are in accordance with results from the previous section where we removed the gaps and analyzed the second-order properties of the patterns only within the homogeneous sub-region. For the pattern shown in Fig. 3A Ripley's L-function indicates regularity at scales $r = 1-3$, and random distribution at scales $r > 3$ (Fig. 5C, 7C). The O-ring statistic indicates regularity at scales $r = 1-2$, a significant aggregation peak at scale $r = 4$ and random distribution at all other scales (Fig. 5B, 7D). For the random pattern with a gap in the center (Fig. 6A) both, the L- and the O-function, reveal the known randomness of the pattern (Fig. 7E, F).

Discussion

In this article, we reviewed current methods in point pattern analysis based on second-order statistics. Over the past 20 years, many useful approaches and methods of analysis utilizing second-order statistics have been explored by statisticians interested in spatial point processes, and point pattern analysis based on Ripley's K has been increasingly used in ecology. However, the full range of methods that is available has not been adopted widely by ecologists, and important problems and pitfalls in their application have not been fully recognized. This might be largely due to the added complexity that is required for implementation of non-standard methods, and due to a lack of appropriate software for their application. Critical issues in point pattern analysis are (1) use of an appropriate method of edge correction, (2) use of specific methods to account for heterogeneity if the pattern is univariate, and (3) selection of an appropriate null model that is used in assessing the observed data, especially for bivariate patterns.

To make methods of second-order statistics more accessible to ecologists we reviewed (1) analytical and numerical methods for implementation of two complementary second order statistics, Ripley's K-function and the O-ring statistic, (2) methods for edge correction, (3) a variety of specific null models for univariate and bivariate patterns, and (4) methods to account for heterogeneity in univariate patterns. Additionally we provide our own software that was implemented following the numerical approach described in the methods section. This software enables ecologists to use most of the standard and non-standard methods reviewed here.

The O-ring statistic

We advocate the use of the O-ring statistic or that of the closely related pair-correlation function, which were both proposed two decades ago by Galiano (1982), and Ripley (1981), respectively, but almost forgotten in later years. We find it quite curious that Ripley's K, which is only one of two potential options of second-order statistics, is widely used while the other option, the pair-correlation function, has rarely been used. One explanation would be that the pair-correlation function $g(r)$ is conceptually too close to spatial correlograms and variograms (Upton and Fingleton 1985). Another explanation is that pair-correlation function has been neglected because it is deterministically related to Ripley's K (Eq. 4): $\lambda g(r)$ is based on the frequency distribution of distances r between all pairs of points while $\lambda K(r)$ is based on the corresponding accumulated frequency distribution.

The scientific question at hand may require the use of the accumulative statistic or the use of the non-accumu-

lative statistic. For example, if the negative effect of competition is hypothesized to work only up to a certain distance, an accumulative statistics may be appropriate. On the other hand, if we ask for "critical scales" in patterns which may e.g. be related to biological processes such as competition, facilitation or seed dispersal we may wish to use a non-cumulative second-order statistic where the result at smaller scales does not bias the result at higher scales. For example, the pattern analyzed in Fig. 3 Fig. 5 is the spatial distribution of *Syagrus yatai* palm trees in the National Park El Palmar in the Argentinean province of Entre Ríos (58° 17' Long. W; 31° 50' Lat. S), a temperate savanna ecosystem (W. Batista and M. Lunazzi, unpubl.). The species *S. yatai* reproduces exclusively by sexually produced seeds, which are then dispersed by gravity or animals. The analysis with the O-function, which measures local neighborhood density at different spatial scales, revealed two critical scales: scales $r = 1$ m with significantly less neighboring trees than expected by a random distribution, and $r = 4$ m with a significantly greater density of trees than expected by a random distribution (Fig. 5B). In contrast, the cumulative L-function shows significant repulsion at scales $r = 1 - 2$ m, and no aggregation at scale $r = 4$ m (Fig. 5C). Thus, Ripley's K- or L-function may actually obscure the existence of critical scales. From the results (Fig. 5B) we may derive the hypothesis that the observed pattern is caused by superposition of the processes of competition and seed dispersal that operate at different spatial scales. Seed-dispersal by gravity may cause a distribution of seeds that is inversely related to the distance from the stem. However, strong competition in the neighborhood of a tree counteracts and causes overall repulsion at these scales. Therefore, the critical scale $r = 4$ m may arise because competition is relatively weak for scales $r \geq 4$ m but seeds are still aggregated at this scale.

The O-ring statistic is a probability density function with the straightforward biological interpretation of a local neighborhood density, which is more intuitive than an accumulative measure (Stoyan and Penttinen 2000). Because it is a scale-dependent density function, the O-ring statistic is only marginally biased by virtual aggregation (caused by larger gaps in the study region that violate the assumption that the pattern is homogeneous). As we have shown, virtual aggregation can be a considerable problem in univariate analysis using Ripley's K. Violation of homogeneity, however, can be detected by visualizing the O-ring statistic. Thus, using the O-ring statistic as complement to Ripley's K may be especially useful in situations where possible violation of homogeneity is not obvious from visual inspection of the pattern.

The number of points of a pattern may constrain the use of the O-ring statistic. Because the accumulative Ripley's K uses all pairs of points that are less than

distance r apart, and the O-ring statistic only all pairs of points a distance r apart, the sample size for calculation of the K-function is considerably larger than that for calculation of the O-ring statistic. The use of too narrow rings or analysis of patterns with very few points will produce jagged plots of the O-ring statistic that make them difficult to interpret. This is less an issue for the accumulative K-function.

The grid approximation

The common analytical approach following Eq. 5 uses all pairs of points to derive an estimator for the K and the O-function, and edge correction is based on geometric formulas that sometimes requires quite complex algorithms and can be computationally intensive. Using an underlying grid simplifies the implementation of second-order statistics considerably. However, one may argue that a grid will reduce the accuracy of $\hat{K}(r)$ and $\hat{O}(r)$ at small scales r because all information below the grain of the grid (the cell size) will be lost. This is true since introducing a grid is equivalent to the introduction of small measurement errors (*sensu* Freeman and Ford 2002). However, all point measurements are subject to measurement error (Freeman and Ford 2002) and if the grid exactly accommodates the measurement error, accuracy is not lost. If the number of points of the pattern is large or if the density of points is low, selection of a small grid size will make the analysis slow and one may wish to select a larger grid size to increase computational speed. Freeman and Ford (2002) investigated in detail which magnitude of measurement errors affects the accuracy of Ripley's L-function, and their results can be used to assess a possible loss of accuracy due to a larger grid size. Addition of measurement error reduced the amplitude of $\hat{L}(r)$ and caused the corresponding maximum value of $\hat{L}(r)$ to move to larger distances. This effect was strongest when the scale of errors approached the scale of the underlying pattern. Because of this, inhibition is more sensitive than clustering, and small clusters are more sensitive than large clusters (Freeman and Ford 2002).

Recommendations

We now synthesize our review into a number of recommendations. Note that point-pattern analysis is a descriptive analysis. Even if a particular null model describes your pattern well, it is not appropriate to conclude that the mechanism behind the null model is the mechanism responsible for your pattern. Other mechanisms may lead to exactly the same pattern. However, point-pattern analysis helps to characterize your pattern and to put forward hypotheses on the underlying mechanisms that should be tested in subse-

quent steps in the field. Therefore, we propose an exploratory step-by-step protocol for second-order analysis. Because there are fundamental differences between univariate and bivariate analysis, we will treat them separately.

Univariate second-order analysis

- 1) Visualize the pattern, define a preliminary study region and plot $\hat{L}(r)$ and $\hat{O}(r)$.
- 2) If the size of your biological objects cannot be neglected (i.e. they are large and do not overlap) you might combine a hard-core null model with the null models suggested in the next steps. In this case $\hat{O}(r)$ will be very low for scales up to the size of the objects.
- 3) If there is no indication for strong aggregation (clearly visible clusters in the pattern or a $\hat{O}(r)$ typical for virtual aggregation) use CSR as the null model for detecting aggregation or inhibition. Virtual aggregation (large scale clustering) is indicated by a constant $\hat{O}(r)$ over a range of scales, and at this range $\hat{O}(r)$ is well above the intensity λ of the pattern (Fig. 3B). Smaller-scale clustering is indicated by a steep linearly increasing $\hat{L}(r)$ at smaller scales (Eq. 18). The cluster size is slightly below the value of r where $\hat{L}(r)$ is maximal.
- 4) If step (3) indicates virtual aggregation (i.e. large clusters) exclude the gaps (or use smaller rectangular sub-regions) and apply CSR only in the sub-region without gaps (or in the smaller plot). Think about a biological explanation for the heterogeneity encountered. Perhaps there are obstacles in the study region, or clear environmental heterogeneity that prevent points from occurring in the gap.
- 5) If there is a biological explanation for the heterogeneity encountered in step (3) (e.g. clear differences in soil), you might map the environmental factor and use this map to obtain an intensity function of a heterogeneous Poisson process. Otherwise, you can use the pattern itself to estimate the non-constant first-order intensity λ using the moving window estimator Eq. 15 for simulation of a heterogeneous Poisson process null model. Alternatively, if there is a surrogate pattern for the environmental heterogeneity (e.g. the locations of a different, more common plant species that is hypothesized to be subject to the same environmental factor), use univariate random labeling as the null model for testing whether your pattern is more (or less) clustered than the control.
- 6) If there is no obvious environmental heterogeneity, your pattern may be a realization of a cluster process. Use $\hat{L}(r)$ to obtain rough (initial) estimates of the parameter ρ and σ of a Neyman-Scott

process and fit the parameters using the methods given in Diggle (1983), Batista and Maguire (1998), and Dixon (2002). Use the estimated parameters ρ and σ to simulate confidence envelopes for the Neyman-Scott process null model. Clearly, there are a number of other point-processes you might fit to your data. However, because of small number of points and noisy data, you might not be able to statistically separate them.

- 7) If there is small-scale regularity and larger scale clustering, the expected L-function for the Neyman-Scott process needs to consider the small-scale regularity because the L-function is accumulative and conserves at larger scales some “memory” on the small-scale regularity (Fig. 4). This can be done analogously to Eq. 20. Alternatively, one may use the pair-correlation function (which has no memory) to fit the unknown parameters ρ and σ .

Bivariate second-order analysis

The bivariate analysis is more complicated than the univariate analysis because there are two basic null models (independence and random labeling) and because null models from the univariate case can be combined in several ways to obtain specific bivariate null models. Therefore, it is especially important to define the biological question, the hypothesis, and the biological circumstances carefully to be able to find an adequate null model.

- 1) Visualize the patterns and perform univariate analysis of both patterns. Define the basic null hypothesis (i.e. independence vs random labeling). Otherwise, if both patterns were probably created by the same stochastic process random labeling is appropriate, whereas independence is appropriate if the patterns might be created by independent processes.
- 2) A common environmental factor affected both patterns in the same way: In this case, the two patterns are heterogeneous and are merged in joint clusters. Under this circumstance, a random labeling null model is appropriate. However, there is a heterogeneous Poisson process null model with the same effect: keep the locations of pattern 1 fixed and randomize pattern 2 according to a heterogeneous Poisson process. An appropriate intensity function can be constructed using a moving window estimate of the joined intensity of pattern 1 and pattern 2, but with a relatively small radius R .
- 3) The two patterns were created by different processes: In this case, you might use the toroidal shift null model, i.e. keeping pattern 1 fixed and shifting

the whole of pattern 2 by treating the study region as a torus. Of course, this works only if you have a rectangular study region.

- 4) The two patterns were created by different processes related to different heterogeneous environmental factors: The appropriate null model for this hypothesis is to keep one pattern fixed and preserve the larger-scale heterogeneity of the other pattern, i.e. use a heterogeneous Poisson process to simulate pattern 2, and vice versa. An appropriate intensity function can be constructed using a moving window estimate of the intensity of pattern 2. The radius R of the moving window decides how closely you mimic the heterogeneity of pattern 2.
- 5) The two processes were linked: An example for this possibility is a clustered distribution of seedlings around adult trees e.g. due to a limited range of seed dispersal. In this case, the locations of trees have to be preserved, and the seedlings can be randomized following a Neyman-Scott process null model where the parents are given through the pattern of adult trees. In this case, only one parameter of the cluster process has to be fit since the intensity ρ is given through the density of pattern 1. Note that a similar effect of clustering of seedlings around trees may arise if both patterns are strongly impacted by the same environmental factor.

Acknowledgements – Funding provided by the UFZ-Centre for Environmental Research, Leipzig, the University of Potsdam and Iowa State University enabled the authors to travel between Germany, and USA for co-operative work. TW thanks the graduate students of his 1999 and 2001 courses “Patrones espaciales en ecología: modelos y análisis”, EPG, Facultad de Agronomía, University Buenos Aires, for stimulating discussions that motivated this article. The authors thank W. Batista and M. Lunazzi for kindly providing the data on the palm tree savanna, S. Higgins, E. Revilla, and especially K. Wiegand for assistance during the development of ideas or for comments on earlier drafts of this manuscript.

References

- Bailey, T. C. and Gatrell, A. C. 1995. Interactive spatial data analysis. – Longman Scientific & Technical.
- Barot, S., Gignoux, J. and Menaut, J.-C. 1999. Demography of a savanna palm tree: Predictions from comprehensive spatial pattern analyses. – *Ecology* 80: 1987–2005.
- Batista, J. L. F. and Maguire, D. A. 1998. Modeling the spatial structure of tropical forests. – *For. Ecol. Manage.* 110: 293–314.
- Besag, J. 1977. Contribution to the discussion of Dr. Ripley's paper. – *J. R. Statist. Soc. B* 39: 193–195.
- Boschdorf, O., Schurr, F. and Schumacher, J. 2000. Spatial patterns of plant association in grazed and ungrazed shrublands in the semi-arid Karoo, South Africa. – *J. Veg. Sci.* 11: 253–258.
- Camarero, J. J., Gutierrez, E. and Fortin, M. J. 2000. Spatial pattern of subalpine forest-alpine grassland ecotones in the Spanish Central Pyrenees. – *For. Ecol. Manage.* 134: 1–16.

- Chen, J. Q. and Bradshaw, G. A. 1999. Forest structure in space: a case study of an old growth spruce-fir forest in Chang-baishan Natural Reserve, PR China. – *For. Ecol. Manage.* 120: 219–233.
- Condit, R., Ashton, P. S., Baker, P. et al. 2000. Spatial patterns in the distribution of tropical tree species. – *Science* 288: 1414–1418.
- Coomes, D. A., Rees, M. and Turnbull, L. 1999. Identifying aggregation and association in fully mapped spatial data. – *Ecology* 80: 554–565.
- Cressie, N. A. C. 1991. *Statistics for spatial data.* – Wiley.
- Dale, M. R. T. 1999. *Spatial pattern analysis in plant ecology.* – Cambridge Univ. Press.
- Dale, M. R. T. and Powell, R. D. 2001. A new method for characterizing point patterns in plant ecology. – *J. Veg. Sci.* 12: 597–608.
- Dale, M. R. T., Dixon, P. M., Fortin, M.-J. et al. 2002. Conceptual and mathematical relationships among methods for spatial analysis. – *Ecography* 25: 558–577.
- Diggle, P. J. 1983. *Statistical analysis of spatial point patterns.* – Academic Press.
- Diggle, P. J. 1985. A kernel method for smoothing point process data. – *J. R. Statist. Soc.: Ser. C (Appl. Statist.)* 34: 138–147.
- Dixon, P. M. 2002. Ripley's K function. – *Encyclopedia Environ.* 3: 1796–1803.
- Dungan, J. L., Citron-Pousty, S., Dale, M. R. T. et al. 2002. A balanced view of scaling in spatial statistical analysis. – *Ecography* 25: 626–640.
- Fehmi, J. S. and Bartolome, J. W. 2001. A grid-based method for sampling and analysing spatially ambiguous plants. – *J. Veg. Sci.* 12: 467–472.
- Freeman, E. A. and Ford, E. D. 2002. Effects of data quality on analysis of ecological pattern using the $\hat{K}(d)$ statistical function. – *Ecology* 83: 35–46.
- Galiano, E. F. 1982. Pattern detection in plant populations through the analysis of plant-to-all-plants distances. – *Vegetatio* 49: 39–43.
- Getis, A. and Franklin, J. 1987. Second-order neighborhood analysis of mapped point patterns. – *Ecology* 68: 474–477.
- Goreaud, F. and Pelissier, R. 1999. On explicit formulas of edge effect correction for Ripley's K-function. – *J. Veg. Sci.* 10: 433–438.
- Grau, H. R. 2000. Regeneration patterns of *Cedrela lilloi* (Meliaceae) in northwestern Argentina subtropical montane forests. – *J. Trop. Ecol.* 16: 227–242.
- Grimm, V., Frank, K., Jeltsch, F. et al. 1996. Pattern-oriented modelling in population ecology. – *Sci. Total Environ.* 183: 151–166.
- Gustafson, E. J. 1998. Quantifying landscape spatial pattern: what is the state of the art? – *Ecosystems* 1: 143–156.
- Haase, P. 1995. Spatial pattern analysis in ecology based on Ripley's K-function: introduction and methods of edge correction. – *J. Veg. Sci.* 6: 575–582.
- Haase, P., Pugnaire, F. I., Clark, S. C. et al. 1996. Spatial patterns in a two-tiered semi-arid shrubland in southeastern Spain. – *J. Veg. Sci.* 7: 527–534.
- Haase, P., Pugnaire, F. I., Clark, S. C. et al. 1997. Spatial pattern in *Anthyllis cytisoides* shrubland on abandoned land in southeastern Spain. – *J. Veg. Sci.* 8: 627–634.
- Haase, P. 2001. Can isotropy vs anisotropy in the spatial association of plant species reveal physical vs biotic facilitation? – *J. Veg. Sci.* 12: 127–136.
- Hanus, M. L., Hann, D. W. and Marshall, D. D. 1998. Reconstructing the spatial pattern of trees from routine stand examination measurements. – *For. Sci.* 44: 125–133.
- He, F. L. and Duncan, R. P. 2000. Density-dependent effects on tree survival in an old-growth Douglas fir forest. – *J. Ecol.* 88: 676–688.
- Jeltsch, F., Moloney, K. A. and Milton, S. J. 1999. Detecting process from snap-shot pattern: lessons from tree spacing in the southern Kalahari. – *Oikos* 85: 451–467.
- Kenkel, N. C. 1988. Pattern of self-thinning in jack pine: testing the random mortality hypothesis. – *Ecology* 69: 1017–1024.
- Klaas, B. A., Moloney, K. A. and Danielson, B. J. 2000. The tempo and mode of gopher mound production in a tallgrass prairie remnant. – *Ecography* 23: 246–256.
- Kuuluvainen, T. and Rouvinen, S. 2000. Post-fire understorey regeneration in boreal *Pinus sylvestris* forest sites with different fire histories. – *J. Veg. Sci.* 11: 801–812.
- Kuusinen, M. and Penttinen, A. 1999. Spatial pattern of the threatened epiphytic bryophyte *Neckera pennata* at two scales in a fragmented boreal forest. – *Ecography* 22: 729–735.
- Larsen, D. R. and Bliss, L. C. 1998. An analysis of structure of tree seedling populations on a Lahar. – *Landscape Ecol.* 13: 307–322.
- Levin, S. A. 1992. The problem of pattern and scale in ecology. – *Ecology* 73: 1943–1967.
- Liebholt, A. M. and Gurevitch, J. 2002. Integrating the statistical analysis of spatial data in ecology. – *Ecography* 25: 553–557.
- Lookingbill, T. R. and Zavala, M. A. 2000. Spatial pattern of *Quercus ilex* and *Quercus pubescens* recruitment in *Pinus halepensis* dominated woodlands. – *J. Veg. Sci.* 11: 607–612.
- Martens, S. N., Breshears, D. D., Meyer, C. W. et al. 1997. Scales of above-ground and below-ground competition in a semi-arid woodland detected from spatial pattern. – *J. Veg. Sci.* 8: 655–664.
- Mast, J. N. and Veblen, T. T. 1999. Tree spatial patterns and stand development along the pine-grassland ecotone in the Colorado Front Range. – *Can. J. For. Res.-Rev. Can. Recherche For.* 29: 575–584.
- Moer, M. 1993. Characterizing spatial patterns of trees using stem-mapped data. – *For. Sci.* 39: 756–775.
- O'Driscoll, R. L. 1998. Description of spatial pattern in seabird distributions along line transects using neighbour K statistics. – *Mar. Ecol.-Progr. Ser.* 165: 81–94.
- Pancer-Koteja, E., Szwagrzyk, J. and Bodziarczyk, J. 1998. Small-scale spatial pattern and size structure of *Rubus hirtus* in a canopy gap. – *J. Veg. Sci.* 9: 755–762.
- Pélissier, R. and Goreaud, F. 2001. A practical approach to the study of spatial structure in simple cases of heterogeneous vegetation. – *J. Veg. Sci.* 12: 99–108.
- Penttinen, A., Stoyan, D. and Henttonen, H. M. 1992. Marked point processes in forest statistics. – *For. Sci.* 38: 806–824.
- Perry, J. N., Winder, L., Holland, J. M. et al. 1999. Red-blue plots for detecting clusters in count data. – *Ecol. Lett.* 2: 106–113.
- Perry, J. N., Liebholt, A. M., Rosenberg, M. S. et al. 2002. Illustrations and guidelines for selecting statistical methods for quantifying spatial patterns in ecological data. – *Ecography* 25: 578–600.
- Podani, J. and Czárán, T. 1997. Individual-centered analysis of mapped point patterns representing multi-species assemblages. – *J. Veg. Sci.* 8: 259–270.
- Prentice, I. C. and Werger, M. J. A. 1985. Clump spacing in a desert dwarf shrub community. – *Vegetatio* 63: 133–139.
- Revilla, E. and Palomares, F. 2002. Spatial organization, group living and ecological correlates in low-density populations of Eurasian badgers, *Meles meles*. – *J. Anim. Ecol.* 71: 497–512.
- Ripley, B. D. 1976. The second-order analysis of stationary point processes. – *J. Appl. Probabil.* 13: 255–266.
- Ripley, B. D. 1977. Modeling spatial patterns. – *J. R. Statist. Soc. B* 39: 172–212.
- Ripley, B. D. 1979. Tests of 'randomness' for spatial point patterns. – *J. R. Statist. Soc. B* 41: 368–374.
- Ripley, B. D. 1981. *Spatial statistics.* – Wiley.
- Ripley, B. D. 1982. Edge effects in spatial stochastic processes. Statistic in theory and practice: essays in honour of Bertil Matérn. – *Swedish Univ. Agric. Sci., Umeå*, pp. 242–262.
- Schooley, R. L. and Wiens, J. A. 2001. Dispersion of kangaroo rat mounds at multiple scales in New Mexico, USA. – *Landscape Ecol.* 16: 267–277.

- Simberloff, D. 1979. Nearest-neighbor assessment of spatial configurations of circles rather than points. – *Ecology* 60: 679–685.
- Sterner, F. J., Ribic, C. A. and Schatz, G. E. 1986. Testing for life history changes in spatial patterns of tropical tree species. – *J. Ecol.* 74: 621–633.
- Stoyan, D. and Stoyan, H. 1994. *Fractals, random shapes and point fields. Methods of geometrical statistics.* – John Wiley & Sons.
- Stoyan, D. and Penttinen, A. 2000. Recent application of point process methods in forest statistics. – *Statist. Sci.* 15: 61–78.
- Turner, M. G. 1989. Landscape ecology: the effect of pattern on process. – *Annu. Rev. Ecol. Syst.* 20: 171–197.
- Upton, G. J. G. and Fingleton, B. 1985. *Spatial data analysis by example: volume 1: point pattern and quantitative data.* – John Wiley & Sons.
- Ward, J. S., Parker, G. R. and Ferrandino, F. J. 1996. Long-term spatial dynamics in an old-growth deciduous forest. – *For. Ecol. Manage.* 83: 189–202.
- Wiegand, T., Moloney, K. A. and Milton, S. J. 1998. Population dynamics, disturbance, and pattern evolution: Identifying the fundamental scales of organization in a model ecosystem. – *Am. Nat.* 152: 321–337.
- Wiegand, T., Moloney, K. A., Naves, J. et al. 1999. Finding the missing link between landscape structure and population dynamics: a spatially explicit perspective. – *Am. Nat.* 154: 605–627.
- Wiegand, K., Jeltsch, F. and Ward, D. 2000. Do spatial effects play a role in the spatial distribution of desert-dwelling *Acacia raddiana*? – *J. Veg. Sci.* 11: 473–484.
- Wiens, J. A., Stenseth, N. C., Van Horne, B. et al. 1993. Ecological mechanisms and landscape ecology. – *Oikos* 66: 369–380.