



Statistical inference for random T-tessellations models. Application to agricultural landscape modeling

Katarzyna Adamczyk-Chauvat¹ · Mouna Kassa² · Julien Papaix³ · Kiên Kiêu¹ ·
Radu S. Stoica⁴

Received: 4 February 2022 / Revised: 2 November 2023 / Accepted: 8 November 2023 /

Published online: 5 April 2024

© The Institute of Statistical Mathematics, Tokyo 2024

Abstract

The Gibbsian T-tessellation models allow the representation of a wide range of spatial patterns. This paper proposes an integrated approach for statistical inference. Model parameters are estimated via Monte Carlo maximum likelihood. The simulations needed for likelihood computation are produced using an adapted Metropolis-Hastings-Green dynamics. In order to reduce the computational costs, a pseudolikelihood estimate is derived and then used for the initialization of the likelihood optimization. Model assessment is based on global envelope tests applied to the set of functional statistics of tessellation. Finally, a real data application is presented. This application analyzes three French agricultural landscapes. The Gibbs T-tessellation models simultaneously provide a morphological and statistical characterization of these data.

Keywords Gibbsian T-tessellation · Monte Carlo Maximum Likelihood estimation · Pseudolikelihood · Global envelope test · Agricultural landscape

1 Introduction

A planar tessellation of a bounded window is called a T-tessellation if all its internal vertices have three adjacent edges, two of them subtending a straight angle. T-tessellations may represent various spatial patterns: cracked soil, burnt wood or

Kiên Kiêu died on 27.02.2017 after an earlier version of this paper had been prepared. The remaining authors take full responsibility for the present version.

✉ Katarzyna Adamczyk-Chauvat
Katarzyna.Adamczyk@inrae.fr

¹ Université Paris-Saclay, INRAE, MaIAGE, 78350 Jouy-en-Josas, France

² INSA, 35700 Rennes, France

³ INRAE, BioSP, 84914 Avignon, France

⁴ Université de Lorraine, CNRS, Institut Elie Cartan de Lorraine, Inria, F-54000 Nancy, France

agricultural landscapes (see Fig. 1). Random T-tessellation models allow for spatial variability of such patterns, generating cells with different sizes and shapes. One of the first random T-tessellation models was proposed by Gilbert (1967) to represent needle-shaped crystals. A variant of the model using rectangular cells was studied by Mackisack and Miles (1996). The STable with respect to Iteration (STIT) tessellation model was proposed by Nagel and Weiss (2005) for modeling random crack networks. STIT tessellations are constructed by successive division of tessellation cells in a particular manner. Another example of a random T-tessellation was proposed by Arak et al. (1993) in their paper on a random graph model. For particular choices of its parameters, the model yields T-tessellation patterns.

In recent years, studies in environmental sciences have attempted to link agro-ecological process dynamics to the features of agricultural landscapes. There is a real need for models that are able to describe and to integrate the morphological and statistical characteristics of agricultural landscapes. The paper by Poggi et al. (2018) presents an overview of landscape modeling challenges in agriculture. Within this general context, the work of Ki  u et al. (2013) introduces the Gibbsian framework for T-tessellations. First, a reference measure is proposed. This measure can be thought of as the most random T-tessellation. Next, with respect to this reference measure, Gibbs probability densities are built. These Gibbs densities make it possible to consider the geometry of each cell of a T-tessellation and to take interactions from cell to cell into account. The paper also describes a Metropolis-Hastings-Green (MHG) dynamics that is able to sample such models, together with the corresponding convergence results.

The present paper completes the previous work with the necessary statistical methodology. Based on the previous MHG dynamics, Monte Carlo maximum likelihood is performed. At the same time, pseudolikelihood parameter estimation for T-tessellations is built. Its role is to be used as the initial condition for the Monte Carlo inference in order to reduce computational costs. Model assessment is done via global envelope tests based on the estimation of the distribution function of the statistics that characterize the geometry and the spatial arrangement of tessellation cells. Finally, three French agricultural landscapes are analyzed with the proposed tools.

The use of tessellation methods for modeling agricultural landscapes was first put forward in Gaucherel (2008) and Le Ber et al. (2009). The landscape representation in these approaches is based on Voronoi tessellation and its generalizations (Okabe et al. (2009)). In particular, the random tessellation models that allow interactions between neighboring cells are of interest in this context. Examples of such models

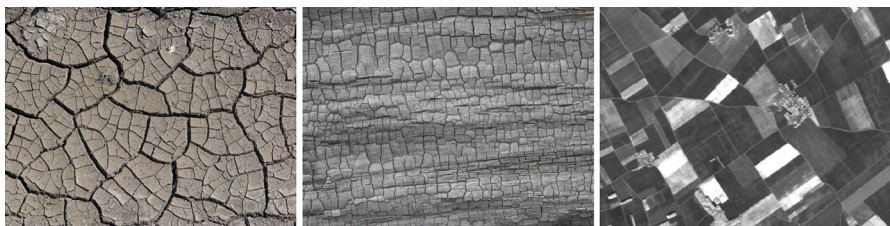


Fig. 1 Examples of spatial patterns that may be represented by a T-tessellation: a cracked soil, a texture of burnt wood, a mosaic of agricultural fields

are Gibbs-Voronoi tessellations (Dereudre (2019), Dereudre and Lavancier (2011)) and Gibbs-Laguerre tessellations, explored in Seidl et al. (2021) for modeling three-dimensional polycrystalline materials. In Le Ber et al. (2009), a Gibbs-Voronoi tessellation was proposed for landscape modeling. The centroids of the cells of an observed agricultural landscape were considered. A Voronoi tessellation was computed for this configuration of points. A Gibbs interaction process was fitted to these data. The interactions take the relative position of the cell centroids into account. The proposed model may be used to differentiate between several types of agricultural landscapes. However, this model does not really characterize the geometry of the observed landscape since such a landscape tessellation is not Voronoi. The Gibbsian T-tessellation model makes it possible to take the spatial variability of landscape patterns into account by direct integration of landscape features into the model definition. The statistical methods proposed in this paper enable direct model fitting to the observed landscapes and allow simulation of tessellations that inherit the observed characteristics of real landscapes.

The plan of this paper is as follows. Section 2 provides the basic definitions and describes the construction of the Gibbs model. The model simulation algorithm is presented in Sect. 3. In Sect. 4, the estimation issue is considered. The maximum likelihood estimation is briefly presented. For the pseudolikelihood approach, introduced for the Gibbsian tessellation model in the unpublished preprint, the construction of the estimate is given in greater detail and includes the description of the theoretical results obtained in Kiêu and Adamczyk-Chauvat (2015), as well as the optimization algorithm of the pseudolikelihood function. Section 5 presents the construction of the global envelope test applied for model assessment. Section 6 continues with the application of the Gibbsian T-tessellation model to the characterization of agricultural landscapes. Three samples of the landscape data are compared. The Gibbs model with potential that highlights the differences between the landscapes is proposed. Model parameters are first estimated by the pseudolikelihood method. Maximum likelihood is used in the second step to refine the pseudolikelihood estimation. The global envelope test is applied to the set of tessellation characteristics. Section 7 presents the main conclusions and highlights some further perspectives. A procedure for approximating a landscape mosaic by a T-tessellation is presented in Appendix A.

2 Gibbs model for T-tessellations

This section starts with a short introduction that provides the basic definitions necessary for model construction. The presentation of tessellations, Gibbs-related models and T-tessellations follows the formalism adopted by Stoyan et al. (1995) and Kiêu et al. (2013), respectively.

2.1 Poisson line process

Let $\mathcal{D}$ be the set of undirected lines in $\mathbb{R}^2$. A line $D(\alpha, p) \in \mathcal{D}$ is characterized by its signed perpendicular distance p from the origin and the angle α between the

line and the x -axis, measured in an anti-clockwise direction and belonging to the interval $[0, \pi)$. A measure δ is defined on the set $\mathcal{D}$ by the following relationship:

$$\int_{\mathcal{D}} \phi(D) \delta(dD) = \frac{1}{\pi} \int_0^\pi \int_{-\infty}^\infty \phi(D(\alpha, p)) d\alpha dp, \quad (1)$$

for any measurable $\phi : \mathcal{D} \rightarrow \mathbb{R}$.

Let $\mathcal{L}$ be the set of all countable subsets of $\mathcal{D}$, referred to as *line configurations*. Let $W \subset \mathbb{R}^2$ be the bounded set of perimeter $b(W)$ and let $\mathcal{D}_W = \mathcal{D} \cap W$.

Definition 1 A homogeneous Poisson line process on W with intensity ν is characterized by the probability distribution λ_W such that:

$$\int_{\mathcal{L}} \phi(L) \lambda_W(dL) = \exp\left(-\frac{b(W)\nu}{\pi}\right) \sum_{n \geq 0} \frac{\nu^n}{n!} \int_{\mathcal{D}_W} \dots \int_{\mathcal{D}_W} \phi(\{D_1 \dots D_n\}) \delta(dD_1) \dots \delta(dD_n) \quad (2)$$

for any measurable $\phi : \mathcal{L} \rightarrow \mathbb{R}$. $\triangleleft$

2.2 T-tessellations

Let $\mathbb{P}$ be the set of all bounded open convex non-empty polygons p in the plane $\mathbb{R}^2$.

Definition 2 A subset of polygons $\Pi \in \mathbb{P}$ is called a *planar tessellation* if the following three conditions are fulfilled:

- (a) $p_1 \cap p_2 = \emptyset$ if p_1 and $p_2 \in \Pi$ and $p_1 \neq p_2$,
- (b) $\bigcup_{p \in \Pi} p^{\text{cl}} = \mathbb{R}^2$, where p^{cl} is a closure of p ,
- (c) if B is a bounded planar set then the set $\{p \in \Pi : p \cap B \neq \emptyset\}$ is finite.

$\triangleleft$

A planar tessellation Π is characterized by the union of all boundaries of polygons in Π , denoted by E_Π . The class $\mathcal{P}$ of all planar tessellations is endowed with the σ -algebra $\sigma(\mathcal{P})$ generated by the sets:

$$\{\Pi \in \mathcal{P} : E_\Pi \cap K\} \quad (3)$$

where K runs through all compact subsets of $\mathbb{R}^2$.

The polygons of the tessellation Π are called *tessellation cells*, and their vertices are *vertices* or *nodes* of Π . A segment lying in E_Π is called a *tessellation edge* if its end-points belong to the set of tessellation vertices and if there is no other vertex belonging to its interior. Two edges are incident to each other if they have

a common end-point. A maximal union of the incident edges belonging to the same line is said to be a *segment* of Π .

In the following, the planar tessellations restricted to a bounded window $W \in \mathbb{R}^2$ will be considered. Tessellation elements that belong to the interior of W will be referred to as *internal elements*.

T-tessellations are a class of planar tessellations characterized by the following property:

Definition 3 A tessellation vertex is said to be a *T-vertex* if it has exactly three incident edges and if two of these edges subtend the angle of π .

A planar tessellation T of a bounded window W is called a *T-tessellation of W* if it satisfies two conditions:

- (i) all the internal vertices of T are T-vertices,
- (ii) for any distinct segments s_1 and s_2 of T , the lines supporting s_1 and s_2 are also distinct.

◁

The space of all T-tessellations of W is denoted by $\mathcal{T}_W$. The σ -algebra $\sigma(\mathcal{T}_W)$ is the restriction of $\sigma(\mathcal{P})$ to the T-tessellations in W . For any T-tessellation $T \in \mathcal{T}_W$, the union E_T of all the boundaries of its cells is equal to the union of the internal segments of T (see Fig. 2). A segment of a T-tessellation consisting of a single edge is called a *non-blocking segment*, whereas a segment with more than one edge is called a *blocking segment*.

2.3 Completely Random T-tessellations and Gibbs models

Random tessellations can be generated in numerous ways (see van Lieshout (2012) for an overview of random tessellation models). The idea behind the construction of the completely random T-tessellation model and its Gibbsian extension is to use the Poisson line process as a skeleton on which a T-tessellation is drawn. Let L_W denote the subset of a line configuration $L \in \mathcal{L}$ containing the lines intersecting W . Let $\mathcal{T}(L_W) \subset \mathcal{T}_W$ denote a set of all T-tessellations of W with segments supported by the lines of L_W (see Fig. 2b).

Definition 4 Let $\{\mathcal{T}_W, \sigma(\mathcal{T}_W)\}$ be the measurable space of T-tessellations of W . The Completely Random T-tessellation Model (CRTT) is given by the following probability distribution on $\{\mathcal{T}_W, \sigma(\mathcal{T}_W)\}$:

$$\mu(A) = \frac{1}{Z} \int_{\mathcal{L}} \left(\sum_{T \in \mathcal{T}(L_W)} \mathbf{1}_A(T) \right) \lambda_W(dL) \text{ for } A \in \sigma(\mathcal{T}_W) \quad (4)$$

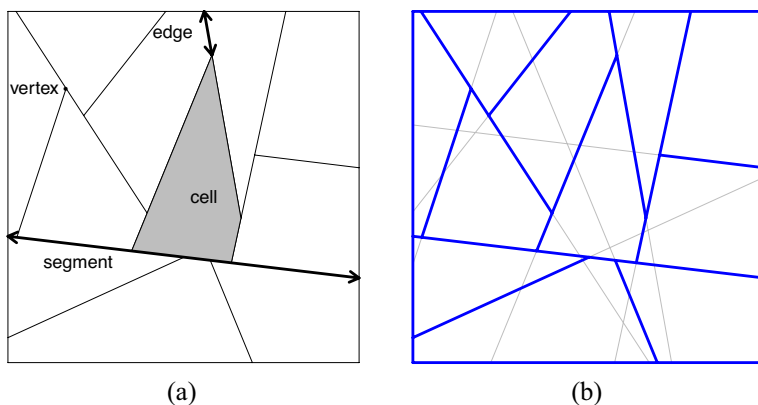


Fig. 2 (a) T-tessellation T of a square window W and the elements of T ; (b) Set of segments of tessellation T (blue lines) together with the supporting line pattern L_W (gray lines); $T \in \mathcal{T}(L_W)$.

where λ_W is a distribution of the Poisson line process with intensity $\nu = 1$, given by the Eq. (2). The symbol $\mathbf{1}_A$ stands for the indicator function. Z is the normalizing constant. $\triangleleft$

The finiteness of the measure μ is stated in Kahn (2015). Figure 3a shows a realization of the CRTT model. The tessellation patterns in an agricultural landscape may exhibit more regularity or alignments. Gibbsian generalizations of the CRTT model make it possible to take such patterns into account.

Definition 5 Let $h : \mathcal{T}_W \rightarrow [0, \infty[$ be the nonnegative function, integrable with respect to the measure μ . The *Gibbs random T-tessellation* of W given by the unnormalized density h is defined as:

$$P(dT) \propto h(T)\mu(dT). \quad (5)$$

$\triangleleft$

The function $U(T) = -\log h(T)$ is called the energy of the model (5).

Hereafter, a parametric family of exponential densities will be considered:

$$h_\theta(T) = \exp(\theta^T t(T)) \quad (6)$$

where $t(T) \in \mathbb{R}^d$ denotes the vector of tessellation statistics and θ belongs to the set of model parameters Θ :

$$\Theta = \{\theta \in \mathbb{R}^d : \int h_\theta(T)\mu(dT) < \infty\}. \quad (7)$$

Example 1 The first example is a model with energy function defined as:

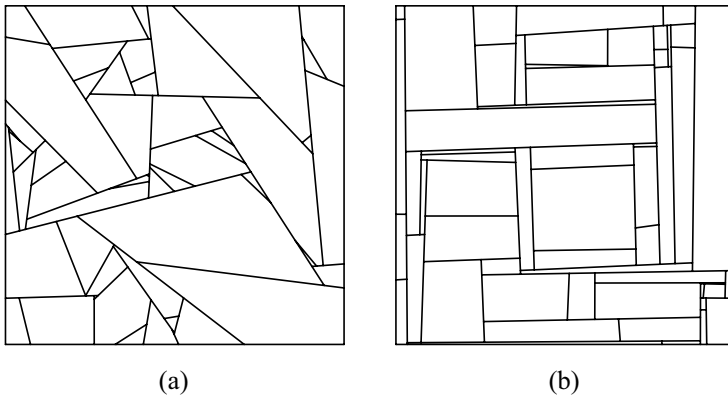


Fig. 3 (a) Simulation of the Gibbs model with $U_\theta(T) = 0.64n_s(T)$ (see Example 1). The source code for generating this example is given in Appendix B.1; (b) simulation of the Gibbs model with $U_\theta(T) = 4.1n_s(T) + 15s_a(T)$ (see Example 2). Both models control tessellation scale, the model simulated in (b) additionally controls the angles between tessellation edges

$$U_\theta(T) = \theta n_s(T) \quad (8)$$

where $n_s(T)$ is the number of internal segments of T . This is a simple generalization of the CRTT model that allows control of the tessellation scale. Consequently, the model (8) will be referred to as the CRTT model with the scale parameter θ (see Fig. 3a). $\triangleleft$

Example 2 Let us consider the Gibbs model given by energy function:

$$U_\theta(T) = \theta s_a(T) = \theta \sum_{c \in C(T)} \sum_{v \in V(c)} \left(\frac{\pi}{2} - \alpha(v) \right) \quad (9)$$

where $C(T)$ is the tessellation cell set and $V(c)$ is the subset of the vertices v of a cell c for which the angle $\alpha(v)$ between the sides incident to v is acute. The contribution of the vertex v tends to 0 if $\alpha(v)$ tends to $\frac{\pi}{2}$. For $\theta > 0$, the model tends to favor tessellations with rectangular cells (see Fig. 3b). $\triangleleft$

3 Simulation algorithm

The Metropolis-Hasting-Green algorithm for the simulation of the Gibbs random T-tessellation described in Kiêu et al. (2013) involves the following steps: the random proposal of type of tessellation update, the random proposal of update of a given type, and the rules for accepting or rejecting the proposals.

Splits, merges and flips Three types of local updates of T-tessellations, illustrated in Fig. 4b are considered:

1. split: a non-blocking segment is added to a tessellation, splitting one of its cells;
2. merge: a non-blocking segment is removed from a tessellation, merging two tessellation cells. A merge can be cancelled by a split and a split can be cancelled by a merge;
3. flip: an edge at the end of a blocking segment is removed while the incident blocked segment is extended. A flip can be cancelled by an inverse flip.

The split and merge operators can be interpreted as the birth of a new segment and the death of the existing one. The restriction to the set of non-blocking segments is required in order to preserve the topology of a T-tessellation. Indeed, removing a blocking segment from a T-tessellation (resp. adding a new one) leads to the occurrence of the vertices that are no longer T-vertices (see Fig. 5a). The flip operator enables a blocking segment to be shortened and reduced to a non-blocking one. In this way, one can progressively remove a blocking segment from a T-tessellation, preserving its topology (see Fig. 5b).

The type of tessellation update is proposed according to a discrete probability distribution:

$$(p_s, p_m, p_f), p_s + p_m + p_f = 1. \quad (10)$$

The spaces of splits, merges and flips that can be applied to a T-tessellation T will be denoted by $\mathbb{S}_T$, $\mathbb{M}_T$ and $\mathbb{F}_T$, respectively.

Uniform measures on $\mathbb{S}_T$, $\mathbb{M}_T$ and $\mathbb{F}_T$ Random proposals of the modification of a given type are defined by the probability distributions constructed on the spaces $\mathbb{S}_T$, $\mathbb{M}_T$ and $\mathbb{F}_T$:

1. A split can be represented by a cell c of T and a line l splitting c . Consequently, $\mathbb{S}_T$ can be identified to the set of pairs (c, l) such that l splits c . The measure on $\mathbb{S}_T$

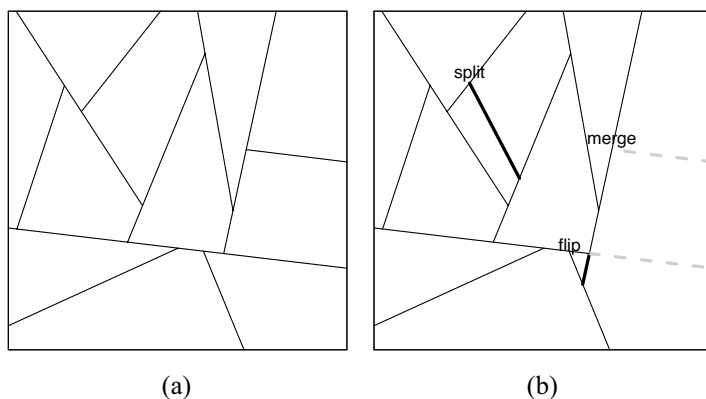


Fig. 4 (a) T-tessellation T from Fig. 2. (b) Three updates of T : split - a non-blocking segment is added (bold line); merge - a non-blocking blocking segment is removed (dashed line); flip - an edge at the end of a blocking segment is removed (dashed line) and the incident segment is extended (bold line)

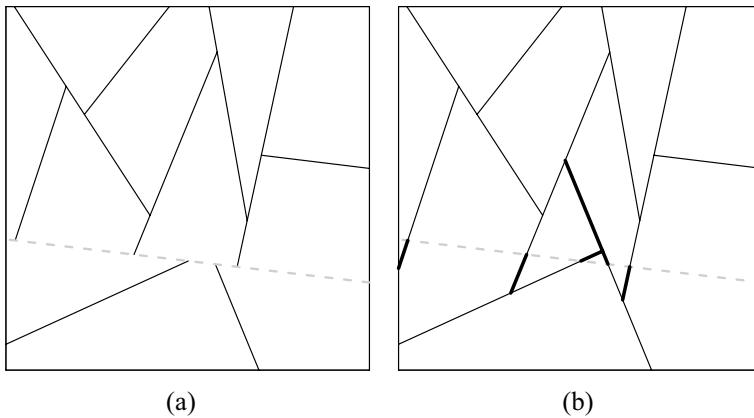


Fig. 5 Left panel: deleting a blocking segment (dashed line) from a T-tessellation. The resulting tessellation is no longer a T-tessellation. Right panel: step-by-step removal of the edges of the blocking segment by a series of flips. The edges added by flips are marked in bold. The resulting tessellation remains a T-tessellation

is defined as a product of a counting measure on the set of tessellation cells and the Haar measure on the space of planar lines. The infinitesimal volume element of this measure can be written as:

$$dS = \sum_{c \in c(T)} \mathbf{1}_{\{c \cap l \neq \emptyset\}} dl \quad (11)$$

where dl is the infinitesimal volume element of the Haar measure. Let $u(T)$ denote the sum of the perimeters of T cells. The uniform probability distribution on $\mathbb{S}_T$ is given by the density:

$$q_s^T(S) = \frac{\pi}{u(T)}, S \in \mathbb{S}_T \quad (12)$$

with respect to the measure (11). Indeed, it can be shown that $\frac{1}{q_s^T(S)}$ is a total mass of the measure (11).

2. A merge can be represented by a non-blocking segment of T to be removed. The set $\mathbb{M}_T$ can be identified with the set of non-blocking segments of T . If $\mathbb{M}_T$ is not empty, a discrete uniform measure can be defined on $\mathbb{M}_T$:

$$q_m^T(M) = \frac{1}{n_{nb}(T)}, M \in \mathbb{M}_T \quad (13)$$

where $n_{nb}(T)$ is the number of non-blocking segments of T .

3. A flip can be represented by a blocking segment of T and by one of its extreme edges to be flipped. If $\mathbb{F}_T$ is not empty, a discrete uniform measure can be defined on $\mathbb{F}_T$:

$$q_f^T(F) = \frac{1}{2n_b(T)}, F \in \mathbb{F}_T \quad (14)$$

where $n_b(T)$ is the number of blocking segments of T . The division by 2 corresponds here to a random choice of the extreme edge of the segment to be flipped.

For a given type t update, an update U is drawn according to the distribution $q_t^T(U)$ and applied to the current tessellation T . The resulting T-tessellation will be denoted by $U(T)$.

Acceptance rules The rules for accepting a new tessellation $U(T)$ are given by the Hastings ratio depending on the proposal distributions and on the energy variation between T and $U(T)$. The Hastings ratio is given by the equation:

$$r_t(T, U) = \frac{h(U(T))}{h(T)} \frac{p_{t^{-1}}}{p_t} \frac{q_{t^{-1}}^{U(T)}(U^{-1})}{q_t^T(U)} \quad (15)$$

where $h(T)$ is the unnormalized density of the model, p_t is given by (10), and $q_t^T(U)$ is one of the uniform measures (12), (13) and (14). The symbol t^{-1} stands for the type of update that cancels the modification introduced by t : $s^{-1} = m$, $m^{-1} = s$ and $f^{-1} = f$. The symbol U^{-1} denotes the inverse of the update U .

The tessellation $U(T)$ is then accepted with the probability:

$$\min\{1, r_t(T, U)\}. \quad (16)$$

The steps of the algorithm are summarized in Algorithm 1.

Algorithm 1 Metropolis-Hastings-Green algorithm for sampling a random Gibbs T-tessellation model (5).

Require: initial tessellation T_0 , density $h(T)$, (p_s, p_m, p_f) , uniform measures $(q_s^T(S), q_m^T(M), q_f^T(F))$.

- 1: The current T-tessellation is T_n .
- 2: Draw the update type t from $\{s, m, f\}$ with probabilities (p_s, p_m, p_f) .
- 3: **if** there is no update of type t applicable to T_n **then**
- 4: $T_{n+1} = T_n$
- 5: **else**
- 6: Draw the update U of type t from the distribution $q_t^{T_n}(U)$.
- 7: Compute the Hastings ratio:

$$r_t(T, U) = \frac{h(U(T))}{h(T)} \frac{p_{t^{-1}}}{p_t} \frac{q_{t^{-1}}^{U(T)}(U^{-1})}{q_t^T(U)}$$

- 8: Accept the tessellation $T_{n+1} = U(T_n)$ with probability

$$\min\{1, r_t(T_n, U)\}.$$

- 9: **end if**
-

Let $\mathbf{T}_n$ be the Markov chain of T-tessellations, defined by Algorithm 1. The properties of $\mathbf{T}_n$, namely its convergence in total variation to the target distribution, were provided in Kiêu et al. (2013). The numerical implementation of the algorithm is available in Adamczyk-Chauvat and Kiêu (2015).

4 Estimation methods

This section presents two estimation methods for the parameters of model (5). The first one is the Monte Carlo Maximum Likelihood given in Penttinen (1984) and Geyer and Thompson (1992) for a general family of distributions specified by the unnormalized densities. Section 4.1 outlines the principles of MCML. The second method is a minimum contrast estimation designed for the Gibbs T-tessellation model in Kiêu and Adamczyk-Chauvat (2015). The extrapolation of the latter to the point process models leads to the definition of a well-known maximum pseudolikelihood estimate. By analogy with point processes, the estimate introduced in Kiêu and Adamczyk-Chauvat (2015) is referred to as the maximum pseudolikelihood estimate (MPL) of the T-tessellation model parameter. Section 4.2 provides a detailed description of the estimate construction, following Kiêu and Adamczyk-Chauvat (2015).

4.1 MCML estimation

The log-likelihood function for the Gibbs model (5) is defined by:

$$l(\theta|T) = -U_\theta(T) - \log Z(\theta) \quad (17)$$

where $U_\theta(T)$ is the energy of the model and $Z(\theta)$ is the unknown normalizing constant. Let $\psi \in \Theta$ be a known parameter value. Maximizing the likelihood (17) is equivalent to maximizing the difference of log-likelihoods:

$$l(\theta|T) - l(\psi|T) = U_\psi(T) - U_\theta(T) - \log \frac{Z(\theta)}{Z(\psi)}. \quad (18)$$

Equation (18) involves the ratio of normalizing constants, which can be expressed as:

$$\frac{Z(\theta)}{Z(\psi)} = \mathbb{E}_\psi \exp (U_\psi(T) - U_\theta(T)) \quad (19)$$

when the expectation in (19) is calculated under the known distribution P_ψ .

Let $\mathbf{T}_n = (T_1, \dots, T_n)$ be a sample from P_ψ , simulated with Metropolis-Hastings-Green algorithm (Algorithm 1). The expected value in (19) can be estimated by its empirical counterpart:

$$\frac{Z(\theta)}{Z(\psi)} \simeq \frac{1}{n} \sum_{i=1}^n \exp (U_\psi(T_i) - U_\theta(T_i)). \quad (20)$$

The estimate of the unknown log-likelihood ratio (18) is obtained by substituting the ratio of normalizing constants with its approximation (20). It is referred to as a Monte Carlo likelihood (MCL) function and denoted by $\hat{l}_n(\theta)$. The maximizer $\hat{\theta}_n$ of $\hat{l}_n(\theta)$ approximates the unknown value of the maximum likelihood estimate $\hat{\theta}$. The accuracy of the Monte Carlo approximation can be assessed due to the asymptotic

normality of $\hat{\theta}_n - \hat{\theta}$ (see Geyer (1994)). Consequently, the standard error of $\hat{\theta}_n - \hat{\theta}$, referred to as the Monte Carlo Standard Error (MCSE), can be controlled by fixing a sufficiently large Monte Carlo sample size.

The likelihood of exponential family models is convex. Hence, the optimization of its Monte Carlo approximation is justified. In practice, the behavior of $\hat{\theta}_n$ depends on the choice of the parameter ψ : the MCL approximation is poor if ψ is far from the true ML estimate. The trust region algorithm (see, for example, Fletcher (1987)) was applied to find the maximum of Monte Carlo likelihood.

Example 3 Let us consider the following Gibbs model:

$$U_{\theta}(T) = -\theta_1 n_s(T) + \theta_2 s_{\angle}(T) \quad (21)$$

where $n_s(T)$ is the number of internal segments of T and $s_{\angle}(T)$ is the angle statistic defined in Example 2. The model was fitted to the tessellation T_0 , simulated with parameters $\theta = (31, 5)$ in the unit square (see Fig. 6a). The MCML estimate was calculated from a Monte Carlo sample of size $n = 400$. The trust region algorithm started from $\psi = (24, 10)$ and converged after 12 iterations at $\hat{\theta}_n = (29.15, 4.76)$ (see Fig. 6b). The distribution of two model statistics calculated over 1000 runs of the fitted model appears to be correctly centered on the observed values (see Fig. 6c). $\triangleleft$

4.2 Pseudolikelihood estimation

The construction of the maximum pseudolikelihood estimate proposed in Kiêu and Adamczyk-Chauvat (2015) is based on the reduced Campbell measures for the Gibbsian T-tessellation model and the Kullback–Leibler divergence extended to nonnegative measures.

Extended Kullback–Leibler divergence

Definition 6 Let ν_1 and ν_2 be two non-negative finite measures on a measurable space $\{\mathcal{E}, \sigma(\mathcal{E})\}$ such that ν_1 is absolutely continuous with respect to ν_2 . The extended Kullback–Leibler divergence of ν_2 from ν_1 is defined as:

$$D_{KL}(\nu_1 || \nu_2) = \int_{\mathcal{E}} \left(\frac{d\nu_1}{d\nu_2}(z) \log \frac{d\nu_1}{d\nu_2}(z) + 1 - \frac{d\nu_1}{d\nu_2}(z) \right) \nu_2(dz) \quad (22)$$

where $\frac{d\nu_1}{d\nu_2}$ denotes the density of ν_1 w.r.t. ν_2 . $\triangleleft$

If ν_1 and ν_2 have the same total mass and, in particular, if ν_1 and ν_2 are probability measures, definition (22) simplifies to:

$$D_{KL}(\nu_1 || \nu_2) = \int_{\mathcal{E}} \left(\frac{d\nu_1}{d\nu_2}(z) \log \frac{d\nu_1}{d\nu_2}(dz) \right) \nu_2(dz). \quad (23)$$

Since the function:

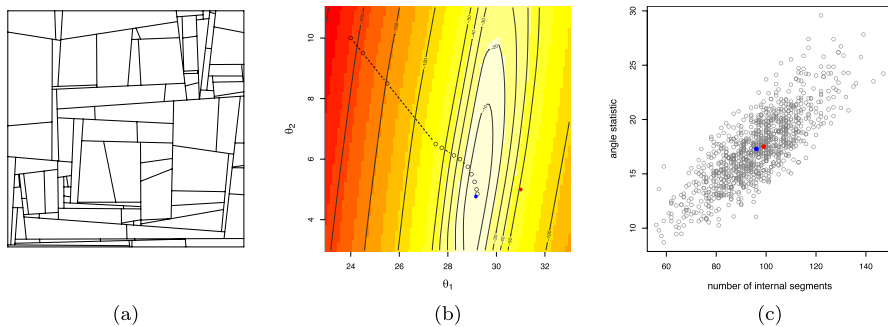


Fig. 6 Estimation of parameters of model (21). **(a)** Training tessellation T_0 simulated with $\theta = (31, 5)$. **(b)** Contour plot of the MCL function. The true parameter is represented by a red dot. The optimization algorithm starts at $\psi = (24, 10)$ and converges after 12 iterations at $\hat{\theta}_n = (29.15, 4.76)$ (blue dot). **(c)** Scatter plot of the values of energy statistics calculated for 1000 simulations of the fitted model. The blue and the red dot represent the mean and the observed values of the statistics, respectively

$$x \rightarrow x \log(x) + 1 - x$$

is non-negative and cancels if and only if $x = 1$, the extended Kullback-Leibler divergence given by Eq. (23) is also nonnegative and cancels only if ν_1 equals ν_2 except from a ν_2 -null subset of $\mathcal{E}$.

Reduced Campbell measures and Papangelou intensities. Let $\mathbf{T}$ be the Gibbs T-tessellation with the distribution $P = h\mu$ on $\{\mathcal{T}_W, \sigma(\mathcal{T}_W)\}$. Let $\mathcal{E}_s = \{(S, T) : T \in \mathcal{T}_W, S \in \mathbb{S}_T\}$ be the product space of $\mathcal{T}_W$ and the splits that can be applied to the tessellations in $\mathcal{T}_W$. Similarly, let $\mathcal{E}_f = \{(F, T) : T \in \mathcal{T}_W, F \in \mathbb{F}_T\}$ be the product space of $\mathcal{T}_W$ and the flips that can be applied to the tessellations in $\mathcal{T}_W$.

Definition 7 The reduced split Campbell measure on $\mathcal{E}_s$ associated with P is defined as:

$$C_s^l(B) = \mathbb{E} \sum_{M \in \mathbb{M}_T} \mathbb{I}_B(M^{-1}, M(\mathbf{T})) \quad \forall B \subset \mathcal{E}_s \quad (24)$$

where M^{-1} denotes the split cancelling the merge M and $M(\mathbf{T})$ denotes the T-tessellation resulting from the application of the merge M to $\mathbf{T}$. $\triangleleft$

Definition 8 The reduced flip Campbell measure on $\mathcal{E}_f$ associated with P is defined as:

$$C_f^l(B) = \mathbb{E} \sum_{F \in \mathbb{F}_T} \mathbb{I}_B(F^{-1}, F(\mathbf{T})) \quad \forall B \subset \mathcal{E}_f \quad (25)$$

where F^{-1} denotes the inverse of the flip F and $F(\mathbf{T})$ denotes the T-tessellation resulting from the application of the flip F to $\mathbf{T}$. $\triangleleft$

Definition 9 The density h of $\mathbf{T}$ is said to be hereditary if and only if:

$$\forall (S, T) \in \mathcal{E}_s \quad h(T) = 0 \Rightarrow h(S(T)) = 0, \quad (26)$$

$$\forall (F, T) \in \mathcal{E}_f \quad h(T) = 0 \Rightarrow h(F(T)) = 0. \quad (27)$$

◁

The following properties of the reduced Campbell measures for the Gibbsian T-tessellations have been derived in Ki  u et al. (2013):

Theorem 1 If $\mathbf{T} \sim h\mu$ and h is assumed to be hereditary, then:

$$C_s^!(ds, dT) = \frac{h(S(T))}{h(T)} dSP(dT) \quad (28)$$

$$C_f^!(df, dT) = \frac{h(F(T))}{h(T)} dFP(dT) \quad (29)$$

where dS is the infinitesimal measure element of the uniform measure defined on $\mathbb{S}_T$ by (12) and dF is the infinitesimal measure element of the counting measure (14) on $\mathbb{F}_T$. ◁

The density $\lambda_s(S, T) = \frac{h(S(T))}{h(T)}$ of the split Campbell measure is called the split Papangelou intensity of $\mathbf{T}$. The flip Papangelou intensity $\lambda_f(F, T)$ of $\mathbf{T}$ is defined in an analogous way.

MPL estimate construction Consider a parametric family of distributions $P_\theta = h_\theta \mu$ on $\{\mathcal{T}_W, \sigma(\mathcal{T}_W)\}$ where $\theta \in \Theta$. The model is assumed to be identifiable:

$$\theta = \theta_0 \iff h_\theta(T) = h_{\theta_0}(T) \quad (30)$$

for μ -almost all $T \in \mathcal{T}_W$. The densities h_θ are assumed to be hereditary. The split and flip Papangelou intensities are denoted by $\lambda_{s,\theta}$ and $\lambda_{f,\theta}$, respectively.

Definition 10 Let T be an observed T-tessellation. The log-pseudolikelihood function of θ given T is defined as:

$$\begin{aligned} \text{LPL}(\theta|T) &= \sum_{M \in \mathbb{M}_T} \log \lambda_{s,\theta}(M^{-1}, M(T)) - \int_{\mathbb{S}_T} \lambda_{s,\theta}(S, T) dS \\ &\quad + \sum_{F \in \mathbb{F}_T} \log \lambda_{f,\theta}(F^{-1}, F(T)) - \int_{\mathbb{F}_T} \lambda_{f,\theta}(F, T) dF. \end{aligned} \quad (31)$$

The maximum pseudolikelihood estimate of θ is defined as:

$$\hat{\theta}_{\text{MPL}} = \arg \max_{\theta \in \Theta} \text{LPL}(\theta|T). \quad (32)$$

◁

The following result states that $-\text{LPL}(\theta|T)$ is a contrast function and that $\hat{\theta}_{\text{MPL}}$ is the associated minimum contrast estimate of θ .

Proposition 1 *If $\mathbf{T} \sim P_{\theta^*}$ then for every $\theta \in \Theta$:*

$$\mathbb{E}_{\theta^*} \text{LPL}(\theta|\mathbf{T}) = M(\theta, \theta^*) = M_s(\theta, \theta^*) + M_f(\theta, \theta^*)$$

where:

$$\begin{aligned} M_s(\theta, \theta^*) &= \int_{\mathcal{E}_s} \log \lambda_{s,\theta}(S, T) C_{s,\theta^*}^\dagger(dS, dT) - \int_{\mathcal{E}_s} \lambda_{s,\theta}(S, T) dSP_{\theta^*}(dT) \\ M_f(\theta, \theta^*) &= \int_{\mathcal{E}_f} \log \lambda_{f,\theta}(F, T) C_{f,\theta^*}^\dagger(dF, dT) - \int_{\mathcal{E}_f} \lambda_{f,\theta}(F, T) dFP_{\theta^*}(dT). \end{aligned}$$

◁

Proof The expected value of the two terms of the log-pseudolikelihood (31) involving splits is equal to:

$$\begin{aligned} & \int_{T_W} \sum_{M \in \mathbb{M}_T} \log \lambda_{s,\theta}(M^{-1}, M(T)) P_{\theta^*}(dT) - \int_{T_W} \int_{\mathbb{S}_T} \lambda_{s,\theta}(S, T) dSP_{\theta^*}(dT) \\ &= \int_{T_W} \int_{\mathbb{S}_T} (\log \lambda_{s,\theta}(S, T)) \lambda_{s,\theta^*}(S, T) dSP_{\theta^*}(dT) - \int_{T_W} \int_{\mathbb{S}_T} \lambda_{s,\theta}(S, T) dSP_{\theta^*}(dT). \end{aligned}$$

The first integral has been transformed according to the Georgii-Nguyen-Zessin formula for the splits (see Kiêu et al. (2013)). The above expression can be rewritten as:

$$\int_{\mathcal{E}_s} \log \lambda_{s,\theta}(S, T) C_{s,\theta^*}^\dagger(dS, dT) - \int_{\mathcal{E}_s} \lambda_{s,\theta^*}(S, T) dSP_{\theta^*}(dT) = M_s(\theta, \theta^*).$$

In a similar manner, it can be shown that:

$$\mathbb{E}_{\theta^*} \left(\sum_{F \in \mathbb{F}_T} \log \lambda_{f,\theta}(F^{-1}, F(T)) - \int_{\mathbb{F}_T} \lambda_{f,\theta}(F, T) dF \right) = M_f(\theta, \theta^*).$$

□

Proposition 2 *If $\mathbf{T} \sim P_{\theta^*}$, then for every $\theta \in \Theta$:*

$$\mathbb{E}_{\theta^*} \text{LPL}(\theta|\mathbf{T}) \leq \mathbb{E}_{\theta^*} \text{LPL}(\theta^*|\mathbf{T}). \quad (33)$$

The equality holds if and only if $\theta = \theta^*$. ◁

Proof (outline)

Consider the difference:

$$M(\theta^*, \theta^*) - M(\theta, \theta^*) = M_s(\theta^*, \theta^*) - M_s(\theta, \theta^*) + M_f(\theta^*, \theta^*) - M_f(\theta, \theta^*).$$

The expression involving the splits can be transformed as follows:

$$\begin{aligned} M_s(\theta^*, \theta^*) - M_s(\theta, \theta^*) &= \\ &= \int_{\mathcal{E}_s} \log \frac{\lambda_{s,\theta^*}}{\lambda_{s,\theta}}(S, T) \lambda_{s,\theta^*}(S, T) dSP_{\theta^*}(dT) + \int_{\mathcal{E}_s} (\lambda_{s,\theta}(S, T) - \lambda_{s,\theta^*}(S, T)) dSP_{\theta^*}(dT) \\ &= \int_{\mathcal{E}_s} \left(\frac{\lambda_{s,\theta^*}}{\lambda_{s,\theta}}(S, T) \log \frac{\lambda_{s,\theta^*}}{\lambda_{s,\theta}}(S, T) + \left(1 - \frac{\lambda_{s,\theta^*}}{\lambda_{s,\theta}}(S, T)\right) \right) \lambda_{s,\theta}(S, T) dSP_{\theta^*}(dT) \\ &= \int_{\mathcal{E}_s} \left(\frac{\lambda_{s,\theta^*}}{\lambda_{s,\theta}}(S, T) \log \frac{\lambda_{s,\theta^*}}{\lambda_{s,\theta}}(S, T) + \left(1 - \frac{\lambda_{s,\theta^*}}{\lambda_{s,\theta}}(S, T)\right) \right) \alpha_{\theta,\theta^*}^s(dS, dT) \end{aligned}$$

where:

$$\alpha_{\theta,\theta^*}^s(dS, dT) = \lambda_{s,\theta}(S, T) dSP_{\theta^*}(dT).$$

Let us notice that the measure $\alpha_{\theta,\theta^*}^s$ dominates the measure $\alpha_{\theta^*,\theta^*}^s$:

$$\alpha_{\theta^*,\theta^*}^s(dS, dT) = \lambda_{s,\theta^*}(S, T) dSP_{\theta^*}(dT) = \frac{\lambda_{s,\theta^*}}{\lambda_{s,\theta}}(S, T) \alpha_{\theta,\theta^*}^s(dS, dT).$$

Consequently:

$$\begin{aligned} M_s(\theta^*, \theta^*) - M_s(\theta, \theta^*) &= \\ &= \int_{\mathcal{E}_s} \left(\frac{d\alpha_{\theta^*,\theta^*}^s}{d\alpha_{\theta,\theta^*}^s}(S, T) \log \frac{d\alpha_{\theta^*,\theta^*}^s}{d\alpha_{\theta,\theta^*}^s}(S, T) + \left(1 - \frac{d\alpha_{\theta^*,\theta^*}^s}{d\alpha_{\theta,\theta^*}^s}(S, T)\right) \right) \alpha_{\theta,\theta^*}^s(dS, dT) \\ &= D_{KL}(\alpha_{\theta^*,\theta^*}^s || \alpha_{\theta,\theta^*}^s). \end{aligned}$$

Similarly, it can be shown that:

$$M_f(\theta^*, \theta^*) - M_f(\theta, \theta^*) = D_{KL}(\alpha_{\theta^*,\theta^*}^f || \alpha_{\theta,\theta^*}^f)$$

where:

$$\alpha_{\theta,\theta^*}^f(dF, dT) = \lambda_{f,\theta}(F, T) dFP_{\theta^*}(dT).$$

It follows that:

$$M(\theta^*, \theta^*) - M(\theta, \theta^*) = D_{KL}(\alpha_{\theta^*,\theta^*}^s || \alpha_{\theta,\theta^*}^s) + D_{KL}(\alpha_{\theta^*,\theta^*}^f || \alpha_{\theta,\theta^*}^f). \quad (34)$$

It follows from the properties of the extended Kullback–Leibler divergence that the difference $M(\theta^*, \theta^*) - M(\theta, \theta^*)$ is non-negative. If $M(\theta^*, \theta^*) - M(\theta, \theta^*) = 0$, then:

$$D_{KL}(\alpha_{\theta^*,\theta^*}^s || \alpha_{\theta,\theta^*}^s) = D_{KL}(\alpha_{\theta^*,\theta^*}^f || \alpha_{\theta,\theta^*}^f) = 0.$$

Consequently $\lambda_{s,\theta}(S, T) = \lambda_{s,\theta^*}(S, T)$ except from the $\alpha_{\theta,\theta^*}^s$ -null set and $\lambda_{f,\theta}(F, T) = \lambda_{f,\theta^*}(F, T)$ except from the $\alpha_{\theta,\theta^*}^f$ -null set. Since the split Papangelou intensity and flip Papangelou intensity uniquely determine the density of $\mathbf{T}$, it follows that $h_\theta = h_{\theta^*}$ except from a μ -null set (see Ki  u and Adamczyk-Chauvat (2015) for the detailed proof). Since the model is assumed to be identifiable, it follows that $\theta = \theta^*$. $\square$

This result shows that $\mathbb{E}_{\theta^*}(-\text{LPL}(\theta|\mathbf{T}))$ has a global unique minimum at $\theta = \theta^*$. Consequently, $\hat{\theta}_{\text{MPL}}$ minimizing $-\text{LPL}(\theta|T)$ is a minimum contrast estimate of θ .

Algorithm for approximating the maximum of pseudolikelihood For the canonical exponential family of distributions (6), the log-pseudolikelihood simplifies to:

$$\begin{aligned} \text{LPL}(\theta|T) = & -\theta^T \sum_{M \in \mathbb{M}_T} \Delta t(M, T) - \int_{\mathbb{S}_T} \exp(\theta^T \Delta t(S, T)) dS \\ & + \theta^T \sum_{f \in \mathbb{F}_T} \Delta t(F, T) - \sum_{f \in \mathbb{F}_T} \exp(\theta^T \Delta t(F, T)) \end{aligned} \quad (35)$$

where:

$$\begin{aligned} \Delta t(M, T) &= t(M(T)) - t(T), \\ \Delta t(S, T) &= t(S(T)) - t(T), \\ \Delta t(F, T) &= t(F(T)) - t(T). \end{aligned}$$

A simple approach for approximating the log-pseudolikelihood consists in discretizing the integral on $\mathbb{S}_T$ that appears in Eq. (35). Let S_D be a finite set of so-called dummy splits of T . The number of elements of S_D is denoted $|S_D|$. The integral in (35) is replaced by a weighted sum:

$$\int_{\mathbb{S}_T} \exp(\theta^T \Delta t(S, T)) dS \approx \frac{u(T)}{\pi} \frac{1}{|S_D|} \sum_{S \in S_D} \exp(\theta^T \Delta t(S, T)). \quad (36)$$

The approximate log-pseudolikelihood function is denoted by $\text{LPL}_d(\theta|T; S_D)$.

A straightforward approximation of the log-pseudolikelihood gradient is given by:

$$\begin{aligned} \nabla \text{LPL}_d(\theta|T; S_D) = & - \sum_{M \in \mathbb{M}_T} \Delta t(M, T) - \frac{u(T)}{\pi |S_D|} \sum_{S \in S_D} \Delta t(S, T) \exp(\theta^T \Delta t(S, T)) - \\ & \sum_{F \in \mathbb{F}_T} \Delta t(F, T) - \sum_{F \in \mathbb{F}_T} \Delta t(F, T) \exp(\theta^T \Delta t(F, T)). \end{aligned} \quad (37)$$

Similarly, the Hessian of the log-pseudolikelihood can be approximated by:

$$H(\text{LPL}_d)(\theta|T; S_D) = -\frac{u(T)}{\pi|S_D|} \sum_{S \in S_D} \Delta t(S, T) \Delta t(S, T)^T \exp(\theta^T \Delta t(S, T)) - \sum_{F \in \mathbb{F}_T} \Delta t(F, T) t(F, T)^T \exp(\theta^T \Delta t(F, T)). \quad (38)$$

The integral on splits in the log-pseudolikelihood can be written as an expectation with respect to the uniform distribution on the space $\mathcal{S}_T$. The estimation criterion consists of deterministic terms and of an expectation with respect to a random variable that can be simulated. Algorithm 2 for approximating the maximum of pseudolikelihood, proposed in Ki  u and Adamczyk-Chauvat (2015), alternates between the updates of the parameter according to a Newton optimization method and the estimation of the expectation.

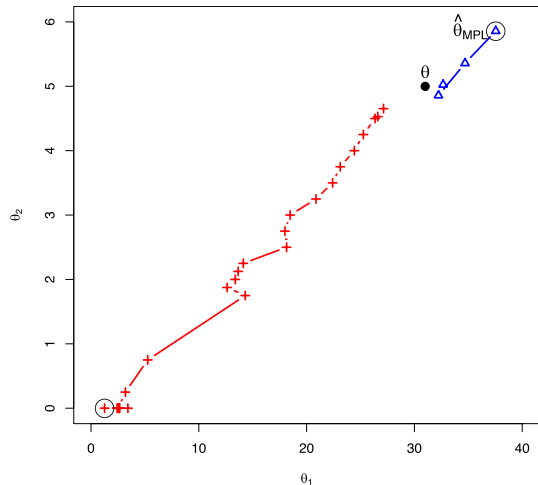
Algorithm 2 Algorithm NOIS (Newton optimization and increasing splitting) for approximating the pseudolikelihood maximum.

Require: a T-tessellation T , the step size ϵ to be used in θ 's update, the tolerance parameter δ

- 1: $m \leftarrow$ number of non-blocking segments of T
 - 2: $\theta \leftarrow$ null vector
 - 3: Compute $\Delta t(M, T)$ for $M \in \mathbb{M}_T$
 - 4: Compute $\Delta t(F, T)$ for $F \in \mathbb{F}_T$
 - 5: $S_D \leftarrow$ sample of m independent and uniform splits S of T
 - 6: Compute the $\Delta t(S, T)$'s
 - 7: **repeat**
 - 8: $G \leftarrow \nabla \text{LPL}_d(\theta|T; S_D)$ as given by (37)
 - 9: $H \leftarrow H(\text{LPL}_d)(\theta|T; S_D)$ as given by (38)
 - 10: $\theta \leftarrow \theta + \epsilon H^{-1}G$
 - 11: Add m independent and uniform splits of T to S_D
 - 12: Update the list of $\Delta t(S, T)$'s where $S \in S_D$
 - 13: $L \leftarrow \text{LPL}_d(\theta|T; S_D)$
 - 14: $\Delta \leftarrow$ variation of L
 - 15: **until** $|\Delta| \leq \delta (|L| + \delta)$.
-

Example 4 The MPL estimate calculated for the model from Example 3 is equal to $\hat{\theta}_{\text{MPL}} = (37.5, 5.8)$. It overestimates both parameters but provides an initial value for MCL maximization that ensures the convergence of the algorithm in only four steps. If no information about the parameters is available, the computation of the MCL estimate may become much slower. For example, one possible strategy would be to first fit a one-parameter model containing the number of internal segment statistic, starting from an arbitrary value of θ_1 , and to initialize the algorithm with $(\hat{\theta}_1, 0)$. Figure 7 illustrates the advantage of MPL initialization compared to such a strategy. Indeed, when applying this strategy, the convergence is reached after 24 iterations and the computation time increases by a factor of 9. $\triangleleft$

Fig. 7 Comparison of two strategies of MCML algorithm initialization for model (21). The first strategy (blue line with triangle symbols) starts from the MPL estimate: the algorithm converges after four iterations. The second strategy (red line with + symbols) fits a one-parameter submodel of (21) starting from an arbitrary value of θ_1 . The value of $(\hat{\theta}_1, 0)$ is used to initialize the MCML algorithm. The convergence is reached in 24 steps and the computation time increases by a factor of 9. The true parameter $\theta = (31, 5)$ is represented by a black dot



5 Goodness of fit of the model

Goodness of fit of the model can be assessed by comparing summary statistics of the observed tessellation pattern with its counterparts calculated from model simulations. The global envelope test is a diagnostic tool that can be applied for this purpose. Originally, the method was considered for point pattern analysis (see Myllymäki et al. (2017)).

The global envelope test is a Monte Carlo goodness of fit test based on the test statistic F . Let $F = (F(r_1), \dots, F(r_d))$ be a statistic, calculated for the set of arguments $\{r_1, \dots, r_d\}$. Let $F_1, F_2, \dots, F_s$ be a sample of d -dimensional vectors of function values calculated for the same set of arguments. F_1 is calculated from the data and $F_2, \dots, F_s$ are calculated from the simulations of the null model. Global envelopes are based on an univariate measure M that orders the functions $F_1, F_2, \dots, F_s$ from the less extreme to the most extreme one (for the sake of simplicity, the order is assumed here to be strict). More precisely:

$$M_i > M_j \iff F_i < F_j$$

where the symbol $<$ stands for less extreme. This means that the smaller the measure M_i is, the more extreme the function F_i will be. Under the hypothesis that the observed pattern comes from a null model, the sample $M_1, \dots, M_s$ is *i.i.d.* The hypothesis is rejected at the significance level α if $M_1 \leq m_\alpha$ where m_α is the empirical α -quantile of $M_1, \dots, M_s$. A Monte Carlo p-value of the test is:

$$p = \frac{1}{s} \sum_{i=1}^s \mathbb{1}(M_i \leq M_1).$$

The test is based on the measures leading to the construction of global envelopes that yield a graphical interpretation of test results. A $100(1 - \alpha)\%$ global envelope is a set of d -dimensional vectors $(F_{\text{low}}^\alpha, F_{\text{upp}}^\alpha)$ such that:

$$M_1 > m_\alpha \iff F_1(r_i) \in (F_{\text{low}}^\alpha(r_i), F_{\text{upp}}^\alpha(r_i)) \text{ for all } r_i \in \{r_1, \dots, r_d\}.$$

Under the null model, the probability that the data function F_1 falls outside the band delimited by F_{low}^α and F_{upp}^α in any of d points is equal to α . Different measures and related global envelopes have been introduced (see Myllymäki and Mrkvička (2020) for the summary of methods). Hereafter, the extreme length test will be applied to the cumulative distribution functions or the density functions of tessellation statistics. The computation of the envelopes relies on the R package GET that implements global envelopes for a general set of d -dimensional vectors in various applications (see Myllymäki and Mrkvička (2020)).

6 Application to the landscape data

In agronomy, finding and quantifying relationships between landscape features and spatial processes is an important research goal. Building parametric stochastic models of the landscape that provide its realistic representation is an effective approach for studying such relationships (see Poggi et al. (2018) and Poggi et al. (2021) for an overview of landscape models and their applications).

The current section presents the application of the Gibbs tessellation model and related statistical methods to landscape modeling. The key idea is to propose the energy function that controls the geometrical features that differentiate three landscape samples. The model is fitted to the landscape data approximated by T-tessellations. Model assessment based on the global envelope test is then carried out.

6.1 Data description

Three landscape datasets were compared: the landscape of Selommes (Centre-Val de Loire), an area of intensive cereal crop cultivation; the landscape of Kervidy (Brittany), a region characterized by mixed crop-livestock farming systems (see Poggi et al. (2020)); and the landscape of the Basse Vallée de la Durance (BVD) in the Provence-Alpes-Côte d'Azur region, specialized in apple orchards. These regions were previously referenced in the study on the influence of landscape features on insect pest dynamics (see Ricci et al. (2009)) and in the modeling of pollen dispersal at the landscape scale (see Devaux et al. (2007)). The datasets are represented in Fig. 8.

The three landscapes were limited to domains encompassing a comparable number of fields: 230 for Selommes, 245 for BVD and 287 for Kervidy. The total area of the domains is equal to 6.2 km^2 for Selommes, 5.8 km^2 for Kervidy and 2.1 km^2 for BVD landscape. At first glance, the landscapes exhibit contrasted patterns: from the most regular landscape of BVD with extremely thin rectangular plots, to the landscape of Kervidy with plots of highly variable sizes and shapes, embedded in the mosaic of woods and farmlands.

The observed landscape patterns are different from the T-tessellation defined by Definition 3. Consequently, the algorithm for approximating landscapes by



Fig. 8 Landscapes of Selommès (a); Kervidy (b); and BVD (c). Light green polygons depict agricultural plots; dark green polygons represent non-agricultural areas (mainly woods, villages and farmlands)

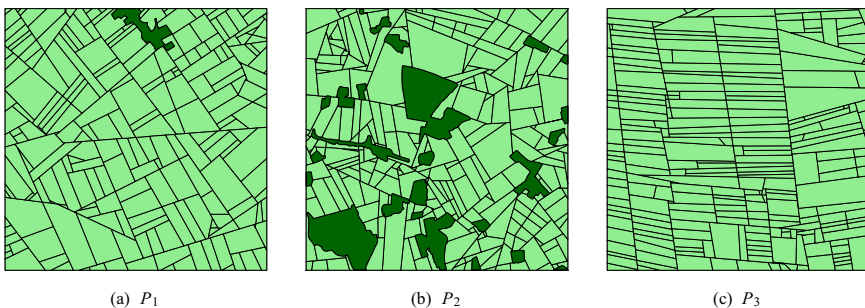


Fig. 9 Landscapes of Selommès, Kervidy and BVD approximated by T-tessellations

T-tessellations was proposed (see Appendix A). The approximation results for the three landscapes are illustrated in Fig. 9.

6.2 Gibbs model for landscape

The landscape model is expected to account for common landscape features (for example, cell shapes close to rectangles), as well as to synthesize the main differences between the three patterns. The Gibbsian model introduced below attempts to combine both aspects.

Let T_1 , T_2 and T_3 be the tessellations approximating the three landscapes (see Fig. 9). A tessellation T_i is assumed to follow a Gibbsian model with the energy defined by:

$$U_{\theta_i}(T_i) = \sum_{j=1}^4 \theta_{i,j} t_j(T_i), i \in \{1, 2, 3\} \quad (39)$$

where:

$t_1(T)$: number of tessellation cells. This term controls the tessellation scale;

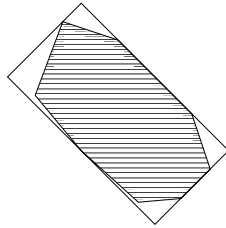


Fig. 10 Smallest rectangle containing a tessellation cell

- $t_2(T)$: sum of squared cell areas. This term measures the departure from tessellations with cells with equal areas, minimizing the value of t_2 ;
- $t_3(T)$: angle statistic introduced in model (9). This term measures the departure from rectangular tessellations;
- $t_4(T)$: number of elongated cells. A measure of cell elongation is defined as the length-to-width ratio of the rectangle with the smallest area containing the cell (see Fig. 10); a cell is said to be elongated if this ratio exceeds a threshold l_0 . In this example $l_0 = 4$, the median of the length-to-width ratio of BVD landscape cells.

The maximum pseudolikelihood estimates of model parameters were calculated in order to initialize the MCML estimate computations. Figure 11 compares two estimation

Table 1 Estimation results for the Gibbs model (39) fitted to the three landscapes: Maximum pseudolikelihood estimate $\hat{\theta}_{\text{MPL}}$, Monte Carlo Maximum Likelihood estimate $\hat{\theta}_n$ ($n = 500$), Standard Error of Maximum Likelihood estimate $\hat{\theta}$ and Monte Carlo Standard Error

Model parameter	Landscape	$\hat{\theta}_{\text{MPL}}$	$\hat{\theta}_n$	$\text{SE}(\hat{\theta})$	MCSE
θ_1 (scale)	P_1	1.62	1.99	0.12	$1.35 \cdot 10^{-6}$
	P_2	1.43	1.70	0.11	$1.05 \cdot 10^{-6}$
	P_3	2.32	2.11	0.13	$2.97 \cdot 10^{-6}$
θ_2 (areas)	P_1	68.36	181.19	66.48	0.17
	P_2	23.73	79.19	56.34	0.20
	P_3	1361.46	2561.46	679.28	28.21
θ_3 (angles)	P_1	0.95	2.23	0.15	$3.28 \cdot 10^{-6}$
	P_2	0.77	1.26	0.12	$1.59 \cdot 10^{-6}$
	P_3	1.72	2.21	0.16	$4.41 \cdot 10^{-6}$
θ_4 (elongation)	P_1	0.48	0.28	0.17	$2.93 \cdot 10^{-6}$
	P_2	0.31	0.07	0.15	$3.44 \cdot 10^{-6}$
	P_3	-0.64	-0.63	0.08	$7.95 \cdot 10^{-6}$

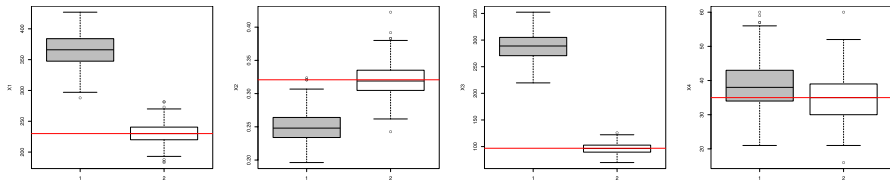


Fig. 11 Distributions of the statistics $t_1(T), \dots, t_4(T)$, calculated for 500 simulations of the model (39) fitted to the Selommes landscape by: (1) MPL method (gray boxplot), (2) MCML method (white boxplot). Red horizontal lines depict the values of the statistics calculated for the Selommes landscape



Fig. 12 Simulations of the Gibbs model (39) fitted to the (a) Selommes, (b) Kervidy, and (c) BVD landscapes

methods on the example of the Selommes landscape, highlighting the bias of MPL estimates corrected by the MCML method. The size of Monte Carlo samples for MCML computation was fixed at $n = 500$ for all the landscapes. A burn-in period was selected after a graphical inspection of the energy function variations. A sampling period was fixed in order to ensure that the autocorrelation level was less than $r_0 = 0.2$ for all the statistics. The values of both MPL and MCML estimates are provided in Table 1.

The model fitted to the BVD landscape (P_3) tends to generate the tessellations close to the patterns with cells of equal size (high positive value of $\hat{\theta}_{n,2}$) and orthogonal segments (positive value of $\hat{\theta}_{n,3}$). The model favors the patterns with a number of long cells greater than in other landscapes and greater than in the reference model (negative value of $\hat{\theta}_{n,4}$). The model fitted to the Kervidy landscape (P_2) is closest to completely random T-tessellation. The estimates of parameters $\theta_{n,2}$, $\theta_{n,3}$ and $\theta_{n,4}$ are smallest in absolute value compared to the other landscapes.

Figure 12 gives the examples of the fitted model simulations for the three landscapes.

6.3 Model assessment

Landscape model assessment begins by checking whether the model fulfills the expected properties. Figure 13 shows the boxplots of model statistics $t_1(T)$, $t_2(T)$, $t_3(T)$ and $t_4(T)$. Each plot juxtaposes the simulated distributions of a given statistic

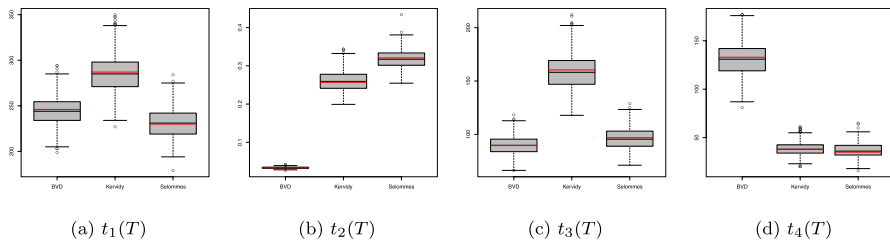


Fig. 13 Model checking: distributions of the statistics $t_1(T)$, $t_2(T)$, $t_3(T)$ and $t_4(T)$, calculated over $n = 500$ runs of the model (39). Red segments represent the observed values of the statistics

in the three landscapes. As expected, all the distributions are correctly centered on the values calculated from landscape data, represented here by red segments. In this sense, the model controls selected features of the landscape. Figure 13 also highlights the differences between the landscapes: the BVD landscape differs from the Selommès and Kervidy landscapes with respect to the statistics $t_2(T)$ and $t_4(T)$. The Kervidy landscape differs from the two others with respect to the statistic $t_3(T)$. Clearly, the model accounts for the observed contrasts.

The goodness of fit of the model with respect to some additional criteria, different from energy statistics, was also assessed. The choice of the statistics was guided by the work of Nagel et al. (2008) and Leon et al. (2020). The following set of features of a typical cell c was considered: the area $a(c)$, the perimeter $p(c)$, the width $w(c)$ and the length $l(c)$ of the smallest rectangle encompassing c and the isoperimetric quotient $r(c)$, also referred to as roundness or sphericity. The latter is defined as the ratio of the area of c to the area of a circle with the same perimeter.

Figure 15 shows the empirical cumulative distribution functions of five statistics for the three landscapes, together with 95% global envelopes based on the extreme length test. The envelopes were constructed from $n = 500$ model runs. The fit of the model to the Selommès landscape is definitely the best: the distribution function of the area, perimeter and width of the cells lies entirely between the envelopes. For the Kervidy landscape, it is the case for the perimeter and the length cumulative distribution function. The model poorly fits the BVD landscape: all the statistics exceed the envelopes in at least one point.

The comparison of the mean values and the coefficients of variance calculated from the observed and simulated data shows that the model tends to keep the mean values of the statistics but overestimates their variability (see Table 2). This conclusion holds for all the statistics except for the cell size $a(c)$. Both the mean and the coefficient of variance of $a(c)$ are close in the observed and model simulations. This is not surprising because the model statistics depend on the first two moments of $a(c)$ distribution. Nevertheless, it is not sufficient to control an entire distribution, at least for the Kervidy and BVD landscapes.

The analysis of spatial pattern of cell centroids completes the model assessment. The centroid (“center of mass”) of a cell is the point whose coordinates are the mean values of coordinates of all points in the cell. The observed centroid patterns are composed from clusters of the aligned points, corresponding to the arrangements of the plots that share a common border (see Fig. 14). The following summary

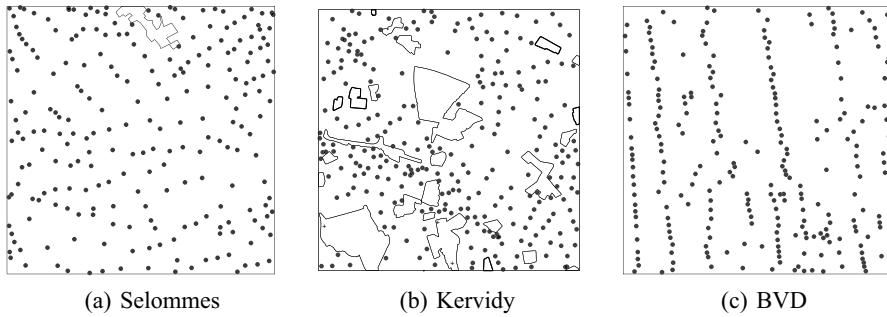


Fig. 14 Centroid patterns of the cells in the observed landscapes

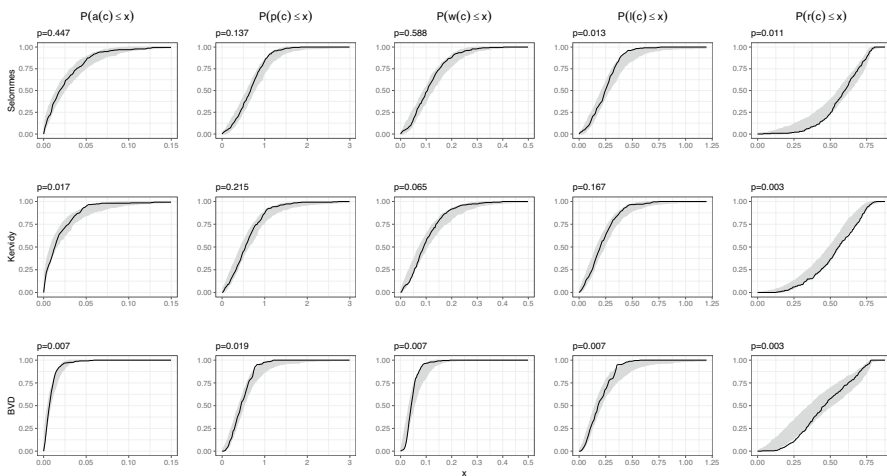


Fig. 15 Empirical cumulative distribution functions (black lines) of tessellation cell characteristics for the three landscapes: size $a(c)$, perimeter $p(c)$, width $w(c)$, length $l(c)$ and roundness $r(c)$. Gray areas represent 95% global envelopes constructed from 500 simulations of the landscape model. P-values of the GET test are given at the top margin of each plot

statistics were considered for comparing the observed and simulated centroid patterns: nearest-neighbor distribution $G(r)$, Besag's L function and empty-space function $F(r)$. The latter refers to the distance $d(u, \mathbf{X})$ from the fixed location $u \in \mathbb{R}^2$ to the nearest point of the stationary point process $\mathbf{X}$, called also empty-space distance. The empty-space function $F(r)$ is a cumulative distribution function of the empty-space distance, namely:

$$F(r) = P(d(u, \mathbf{X}) \leq r)$$

for all distances $r \geq 0$. The calculation of summary statistics was done in spatstat package (see Baddeley et al. (2016)). The estimates of summary statistics and the 95% global envelopes are represented in Fig. 16. The p-values of the GET test are

Table 2 Comparison of landscape data and model simulations. Mean values and coefficients of variations (cv) were calculated for each landscape and compared with their simulated counterparts (averages over $n = 500$ simulations are shown)

Landscape		Tessellation cell statistics									
		$a(c)$		$p(c)$		$w(c)$		$l(c)$		$r(c)$	
		mean	cv	mean	cv	mean	cv	mean	cv	mean	cv
Selommès	Data	0.026	0.97	0.667	0.51	0.113	0.61	0.249	0.51	0.584	0.23
	Model	0.026	0.96	0.700	0.56	0.118	0.63	0.266	0.62	0.563	0.30
Kervidy	Data	0.019	1.23	0.600	0.66	0.103	0.67	0.220	0.63	0.544	0.29
	Model	0.019	1.23	0.613	0.76	0.105	0.77	0.220	0.71	0.507	0.33
BVD	Data	0.008	0.92	0.467	0.49	0.046	0.54	0.198	0.55	0.471	0.37
	Model	0.008	0.92	0.503	0.64	0.049	0.69	0.216	0.74	0.435	0.49

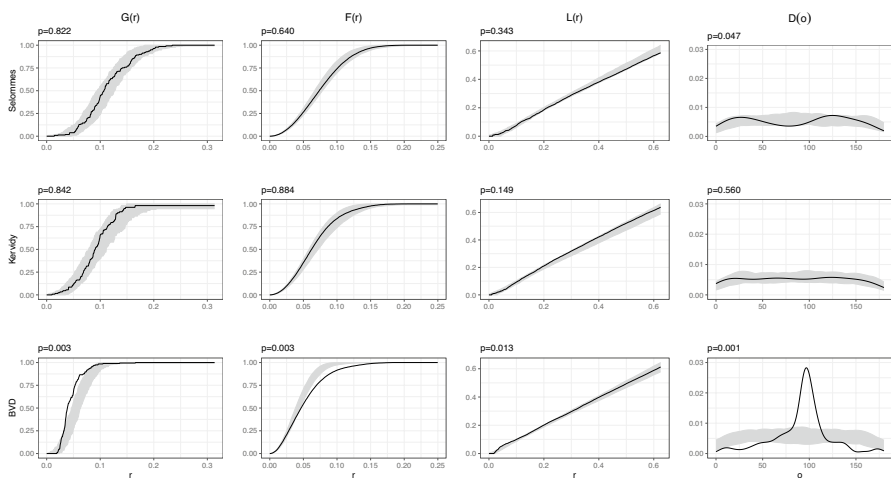


Fig. 16 Estimates of summary statistics (black lines) calculated for cell centroids of the three landscapes: nearest-neighbor function $G(r)$, empty-space function $F(r)$, Ripley function $L(r)$ and nearest-neighbor orientation density $D(o)$. Gray areas represent 95% global envelopes constructed from 500 simulations of the landscape model

higher than 0.05 for all the functions calculated for the Selommès and Kervidy landscapes. The model particularly well fits both nearest-neighbor distribution functions. As for the BVD landscape, the centroid pattern is significantly more aggregated than the simulated patterns.

In addition, the distribution of the direction of the vector joining each centroid to its nearest neighbor was calculated in order to detect anisotropy (see Baddeley et al. (2016)). The estimated density $D(o)$ of nearest-neighbor orientation for the three landscapes is shown in Fig. 16 (the last column of the plots). For the landscapes of Selommès and BVD, the shape of the density function is

determined by the preferential directions of point alignments: bimodal distribution for Selommès and unimodal, extremely peaked distribution for BVD. In the latter, the alignments of points are long and follow one leading direction. In the Selommès landscape, the alignments are shorter and concentrated around two orthogonal directions. In both cases, the estimated density exceeds global envelopes. The centroid pattern in the Kervidy landscape does not exhibit anisotropy, and the density of nearest-neighbor orientation for this landscape lies between the envelopes.

7 Concluding remarks

This paper proposes the extension of statistical inference tools for Gibbsian models to T-tessellations. An application to agricultural landscape modeling is then developed for three landscape samples. The mathematical development presented here is a continuation of the work of Kiêu et al. (2013) and Kiêu and Adamczyk-Chauvat (2015). The main results of this work are outlined in order to provide the reader with a self-contained paper.

The pseudolikelihood approach designed in Kiêu and Adamczyk-Chauvat (2015) for Gibbsian T-tessellations is recommended for model parameter estimation for practical reasons. The main advantage of the pseudolikelihood procedure is the low computational cost of the criterion optimization. For the CRTT model with scale parameter, the analytical form of the estimate is available and may be applied as an initial value in more complex models. On the other hand, even if the bias and dispersion of the estimate decrease as the number of cells in the observed tessellation increases, they may be important for small-scale tessellations. Therefore, the maximum likelihood should be used in the second step to refine the pseudolikelihood estimation. Other inference methods should also be investigated for Gibbsian T-tessellations. For example, Bayesian approaches could be of interest (see Stoica et al. (2017) and Vihrs et al. (2021)).

The model can be simulated using the Metropolis-Hasting-Green algorithm based on three local modifications of a T-tessellation. It is worth noting that for some settings, model simulation becomes prohibitively time-consuming. This is especially the case of the model fitted to the BVD landscape that requires an extremely high burn-in and sampling period. One possible explanation for this behavior is that the acceptance probabilities of the algorithm became small in the presence of the high parameter values estimated for this landscape. A similar problem has already been addressed in Seidl et al. (2021). Alternative simulation algorithms should therefore be investigated. In particular, it seems interesting to test if the algorithm for simulating polygonal Markov fields proposed in Schreiber and van Lieshout (2010) would be adaptable to the T-tessellation Gibbs model.

A multiple testing procedure is recommended in this paper for model assessment. A global envelope test is applied to the set of functional statistics of tessellation. The resulting p-values indicate to what extent the model fits the data and provide some guidelines for model improvement. The choice of benchmark statistics, based here on bibliographic considerations, can be viewed as a starting proposition that does

not account for the correlations between statistics. Building a set of generic tessellation indices for model validation is a question that requires further insight.

The motivation of using Gibbsian T-tessellations is to statistically and morphologically characterize agricultural landscapes. The model with the energy function that accounts for the selected landscape features was proposed and fitted to the landscape data. This is for the first time that the model has been used to analyze this kind of data. The results confirm that it controls the geometrical properties of the landscapes summarized in four model statistics. Model assessment with respect to the additional criteria was carried out using global envelope tests. The results of model fit for the landscapes of Selommès and Kervidy depend on the considered statistic. In general, the model correctly matches the spatial pattern of cell centroids for both landscapes. The goodness of fit of tests related to geometrical features of tessellation cells is more variable. As for the BVD landscape, the model does not correctly handle the high anisotropy and the strong regularity of the observed pattern. The fit of the model is uniformly wrong for this landscape, except from energy statistics. The generalization of the model handling anisotropy could improve the goodness of fit. Furthermore, the energy of the landscape model is composed from the first-order terms: considering higher-order potentials could provide more suitable model. The examples of such potential functions can be found in the paper of Seitz et al. (2021) on the statistical reconstruction of the polycrystalline microstructures using Gibbs-Laguerre tessellations.

Another point should be raised when discussing the goodness of fit. The Gibbs T-tessellation model is based on the assumption that tessellation pattern is driven by local tessellation features. This assumption seems reasonable for the Selommès and the Kervidy landscapes but probably does not hold for the BVD landscape. This landscape is located in the region subject to the action of the Mistral wind. This area is covered by a dense network of hedges, protecting crops from the Mistral. The hedges, delimiting agricultural plots, run orthogonally to the Mistral direction. The spatial structure of the BVD landscape is therefore driven by a hidden variable: presence of the hedges on a much larger spatial scale. For this particular landscape, it could be more relevant to model hedge pattern rather than agricultural plot pattern. A fiber process model with highly correlated fibers could be considered for this purpose.

The modeling framework, illustrated here on the example of an agricultural landscape, can be extended to other applications dealing with datasets that can be approximated by T-tessellations. The goal is to provide clues to the practitioners to help them build the models that match the desired properties of such datasets that are important for the applications.

Acknowledgements This work was supported partly by the French PIA project “Lorraine Université d’Excellence” (ANR-15-IDEX-04-LUE) and by the SMACH Metaprogram (Sustainable Management of Crop Health) of the French National Research Institute for Agriculture, Food and Environment. The BVD and the Kervidy landscape data were kindly provided by Claire Lavigne (INRAE, Centre de PACA) and Sylvain Poggi (INRAE, Centre de Bretagne-Normandie). We are grateful to the authors of the online library “Texturelib” for the examples of tessellation patterns shown in Fig. 1. We would like to thank Ms. Gail Wagman for her contribution to improving the English of the manuscript. We are grateful to the Reviewers for helpful and constructive comments.

A. Algorithm for approximating a landscape pattern by a T-tessellation

The algorithm first approximates a landscape pattern by a line-based tessellation containing the L-vertices, the I-vertices and the X-vertices (see Figs. 8 and 17). In the second step, these vertices are converted into T-vertices.

The algorithm requires the numerical representation of the landscape composed from the list of the plots coordinates, the list of the holes coordinates (the holes correspond to non-arable areas) and the coordinates of the outer domain.

Grouping polygon sides

The algorithm starts by grouping the sides of polygons that are close and nearly aligned. This is done by hierarchical clustering based on the single linkage method. A measure of dissimilarity between the sides can be adjusted to account for specific landscape features. Hereafter, the following dissimilarity matrix was applied:

$$d(c_i, c_j) = \frac{A(C(c_i, c_j))}{l(c_i)l(c_j)} + \frac{2d_{\min}(c_i, c_j)}{l(c_i) + l(c_j)}, \quad (40)$$

where $A(C(c_i, c_j))$ denotes the area of the convex envelope of the sides c_i and c_j , $d_{\min}(c_i, c_j)$ is the smallest Euclidean distance between c_i and c_j , and $l(\cdot)$ denotes the length of the side.

Choice of representative segments

The clusters of sides are then replaced by representative segments. For a given cluster, a regression line is fitted to the endpoints of its sides. A representative segment is defined as the minimal segment containing the projections of all endpoints on the regression line (see Fig. 18). The segments representing the side clusters are inserted into a landscape domain. The resulting segment-based tessellation contains three types of vertices: I-vertices with one incident edge, L-vertices with two incident edges, and X-vertices formed by the crossing segments (see Fig. 17).

Transformation into T-tessellation

In order to transform an I vertex into a T-vertex, the incident edge is either removed or lengthened until it meets another edge. The modification that induces the smallest variation in total internal length of tessellation is selected. An L-vertex is transformed into a T-vertex by lengthening one of the incident edges. The choice of the edge is also based on the comparison of length variation.

The algorithm starts by iteratively removing the I-vertices: the vertex whose suppression induces the smallest variation in length is selected at each iteration. The L-vertices are then removed following the same rule. The resulting tessellation contains the T-vertices and the X-vertices. An X-vertex can be suppressed by moving the endpoint of one of its four incident edges along a neighboring edge. The procedure is repeated until there is no X-vertex left in the tessellation. Figure 17 illustrates the modifications of I-, L- and X-vertices.

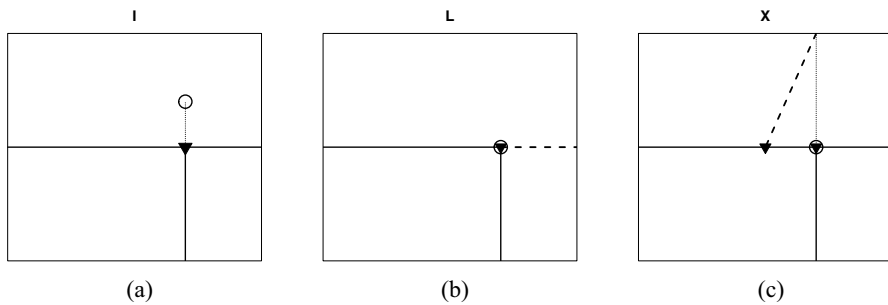
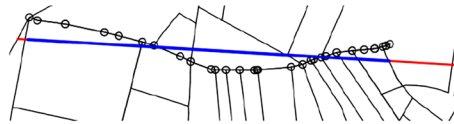


Fig. 17 Transformation of segment-based tessellation into a T-tessellation: deletion of (a) I-vertex; (b) L-vertex; and (c) X-vertex. The original vertices are marked by dots, while the new T-vertices are marked by triangles. Dotted lines depict the removed edges and dashed lines depict the inserted edges

Fig. 18 A cluster of polygon sides replaced by a representative segment (in blue)



B. Numerical examples of T-tessellation construction

B.1 Sampling from the CRTT model

This section gives an example of source code for sampling from the CRTT model with scale parameter θ (see Example 1) with Algorithm 1 implemented in the RLiTe package (see Adamczyk-Chauvat and Kiêu (2015)):

```
require(RLiTe)

tau=1.9                                     # fixing scale parameter (theta=log(tau) in Ex.1)
myTes <- new(TTessel)                       # myTes is defined for T-tessellation
myTes$setDomain(1,1)                       # the domain of myTes is specified
myMod <- ModelCRTT(tau=tau)                 # the CRTT model object is defined ...
myMod$set_ttessel(myTes)                   # ...and assigned to myTes object
mySim <- new(SMFChain,myMod,0.33,0.33)      # initialization of the algorithm
setLiTeSeed(5)                             # random seed
mySim$step(1000)                           # first 1000 steps of the algorithm
plot(myTes)                                # sampled tessellation
```

B.2 Random splits

In this section, the algorithm for simulating the T-tessellation of a bounded window W by a sequence of random splits is provided. It can be applied to create a tessellation of an empty window or to modify the non-empty tessellation, e.g., to add dummy splits to the existing tessellation in the pseudolikelihood optimization algorithm.

The following notation is used:

- $e = (p_1, p_2)$: an oriented edge of a tessellation T (halfedge for short) with the starting point p_1 and the ending point p_2 ;
- $c(e)$: tessellation cell bounded by e and located on the left side of e
- $\mathcal{E}(T)$: ordered list of halfedges of T

Two auxiliary algorithms, 3 and 4, are first introduced and applied in the main algorithm 5.

Algorithm 3 Draw at random a half-edge bounding an internal face of T

DrawAtRandom

Require: $\mathcal{E}(T)$

```

1: found  $\leftarrow$  FALSE
2:  $n \leftarrow$  number of half-edges in  $\mathcal{E}(T)$ 
3: while (! found) do
4:    $ri \leftarrow$  random integer sampled from the uniform distribution on  $\{1, \dots, n\}$ 
5:    $e \leftarrow$   $ri^{\text{th}}$  element of  $\mathcal{E}(T)$ 
6:   if  $c(e)$  internal then found  $\leftarrow$  TRUE
7:   end if
8: end while
9: return  $e$ 
```

Algorithm 4 Make a split starting from the selected halfedge of T

MakeASplit

Require:

$\mathcal{E}(T)$

e_1 : half-edge bounding the cell $c(e_1)$ where the splitting segment starts;

$s \in (0, 1)$: relative position along the half-edge defining where the splitting segments starts.

$a \in (0, \pi)$: angle between the half-edge and the splitting segment.

```

1: Compute the starting point  $p_1(e_1, s)$  of the splitting segment
2: Compute the line  $l(p_1, a)$  containing the splitting segment,
3: Find another half-edge  $e_2 \neq e_1$  bounding the cell  $c(e_1)$  intersecting the line  $l$ 
4: Compute the intersection point  $p_2$  between  $l$  and  $e_2$ 
5: Add the twin half-edges  $(p_1, p_2)$  and  $(p_2, p_1)$  to  $\mathcal{E}(T)$ 
```

Algorithm 5 Random splits

Require:

$\mathcal{E}(T)$

n : number of splits

```

1: for  $i \leftarrow 1$  to  $n$  do
2:    $e \leftarrow$  random half-edge generated by DrawAtRandom( $\mathcal{E}(T)$ )
3:    $s \leftarrow$  random position of the splitting segment,  $s \sim \mathcal{U}(0, 1)$ 
4:    $a \leftarrow$  random direction of the splitting segment,  $a \sim \mathcal{U}(0, \pi)$ 
5:   MakeASplit( $\mathcal{E}(T), e, s, a$ )
6: end for
```

References

- Adamczyk-Chauvat, K., Kiêu, K. (2015). LiTe. <http://kien-kieu.github.io/lite>
- Arak, T., Clifford, P., Surgailis, D. (1993). Point-based polygonal models for random graphs. *Advances in Applied Probability*, 25, 348–372.
- Baddeley, A., Rubak, E., Turner, R. (2016). *Spatial Point Patterns: Methodology and Applications with R*. Boca Raton: Chapman and Hall/CRC Press.
- Dereudre, D. (2019). Introduction to the theory of Gibbs point processes. In D. Coupier (Ed.), *Stochastic Geometry*. Cham: Springer International Publishing.
- Dereudre, D., Lavancier, F. (2011). Practical simulation and estimation for Gibbs Delaunay-Voronoi tessellations with geometric hardcore interaction. *Computational Statistics and Data Analysis*, 55, 498–519.
- Devaux, C., Lavigne, C., Austerlitz, F., Klein, E. F. (2007). Modelling and estimating pollen movement in oilseed rape (*Brassica napus*) at the landscape scale using genetic markers. *Molecular Ecology*, 16, 487–499.
- Fletcher, R. (1987). *Practical Methods of Optimization* (2nd ed.). Chichester: Wiley.
- Gaucherel, C. (2008). Neutral models for polygonal landscapes with linear networks. *Ecological Modelling*, 219, 39–49.
- Geyer, C. J. (1994). On the convergence of Monte Carlo Maximum Likelihood calculations. *Journal of the Royal Statistical Society, Series B*, 56, 261–274.
- Geyer, C. J., Thompson, E. A. (1992). Constrained Monte Carlo Maximum Likelihood for dependent data. *Journal of the Royal Statistical Society, Series B*, 54, 657–699.
- Gilbert, E. N. (1967). Random plane network and needle-shaped crystals. In B. Noble (Ed.), *Applications of Undergraduate Mathematics in Engineering*. New York: Macmillan.
- Kahn, J. (2015). How many T-tessellations on k lines? Existence of associated Gibbs measures on bounded convex domains. *Random Structures & Algorithms*, 47, 561–587.
- Kiêu, K., Adamczyk-Chauvat, K. (2015) Pseudolikelihood inference for Gibbsian T-tessellations. . . and point processes. <https://hal.archives-ouvertes.fr/hal-02793013/>
- Kiêu, K., Adamczyk-Chauvat, K., Monod, H., Stoica, R. (2013). A completely random T-tessellation model and Gibbsian extensions. *Spatial Statistics*, 6, 118–138.
- Le Ber, F., Lavigne, C., Adamczyk, K., Angevin, F., Colbach, N., Mari, J. F., Monod, H. (2009). Neutral modelling of agricultural landscapes by tessellation methods-application for gene flow simulation. *Ecological Modelling*, 220, 3536–3545.
- Leon, R., Nagel, W., Ohser, J., Arscott, S. (2020). Modeling crack pattern by modified STIT tessellations. *Image Analysis and Stereology*, 39, 33–46.
- Mackisack, M., Miles, R. (1996). Homogeneous rectangular tessellations. *Advances in Applied Probability*, 28, 993–1013.
- Myllymäki, M., Mrkvička, T. (2020). GET: Global envelopes in R. [arXiv:1911.06583](https://arxiv.org/abs/1911.06583) [statME] [arXiv:1911.06583](https://arxiv.org/abs/1911.06583)
- Myllymäki, M., Mrkvička, T., Grabarnik, P., Seijo, H., Hahn, U. (2017). Global envelope tests for spatial processes. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 79, 381–404. <https://doi.org/10.1111/rssb.12172>
- Nagel, W., Weiss, V. (2005). Crack STIT tessellations: Characterization of stationary random tessellations stable with respect to iteration. *Advances in Applied Probability*, 37, 859–883.
- Nagel, W., Mecke, J., Ohser, J., Weiss, V. (2008). A tessellation model for crack patterns on surfaces. *Image Analysis and Stereology*, 27, 73–78.
- Okabe, A., Boots, B., Sugihara, K., Chiu, S. N. (2009). *Spatial Tessellations: Concepts and Applications of Voronoi Diagrams*. New York: Wiley.
- Penttinen, A. (1984). In: *Modelling Interaction in Spatial Point Patterns: Parameter Estimation by the Maximum Likelihood Method*, vol 7, Jyväskylä Studies in Computer Science, Economics, and Statistics. University of Jyväskylä
- Poggi, S., Papaix, J., Lavigne, C., Angevin, F., Le Ber, F., Parisey, N., Ricci, B., Vinatier, F., Wohlfahrt, J. (2018). Issues and challenges in landscape models for agriculture: from the representation of agroecosystems to the design of management strategies. *Landscape Ecology*, 33, 1679–1690.
- Poggi, S., Delaplace, A., Pichelin, P., Le Cointe, R. (2020). History of land uses over the period 2013–2019 in the Kervidy-Naizin watershed (Brittany, France). <https://doi.org/10.15454/ATYOBO>

- Poggi, S., Vinatier, F., Hannachi, M., Sanz Sanz, E., Rudi, G., Zamberletti, P., Tixier, P., Papaix, J. (2021). How can models foster the transition towards future agricultural landscapes? *Advanced in Ecological Research*, 64, 305–368.
- Ricci, B., Franck, P., Toubon, J. F., Bouvier, J. C., Sauphanor, B., Lavigne, C. (2009). The influence of landscape on insect pest dynamics: a case study in southeastern France. *Landscape Ecology*, 24, 337–349.
- Schreiber, T., van Lieshout, M. N. M. (2010). Disagreement loop and path creation/annihilation algorithms for binary planar Markov fields with applications to image segmentation. *Scandinavian Journal of Statistics*, 37, 264–285.
- Seitl, F., Petrich, L., Staněk, J., Krill, C., III., Schmidt, V., Beneš, V. (2021). Exploration of Gibbs-Laguerre tessellations for three dimensional stochastic modeling. *Methodology and Computing in Applied Probability*, 23, 669–693.
- Stoica, R. S., Philippe, A., Gregori, P., Mateu, J. (2017). An ABC method for posterior sampling of marked point processes. *Statistics and Computing*, 27, 1225–1238.
- Stoyan, D., Kendall, W. S., Mecke, J. (1995). *Stochastic geometry and its applications* (2nd ed.). Chichester: Wiley.
- van Lieshout, M. N. M. (2012). An introduction to planar random tessellation models. *Spatial Statistics*, 1, 40–49.
- Vihrs, N., Möller, J., Gelfand, A. E. (2021). Approximate Bayesian inference for a spatial point process model exhibiting regularity and random aggregation. *Scandinavian Journal of Statistics*, 48, 969–1000.

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.

Springer Nature or its licensor (e.g. a society or other partner) holds exclusive rights to this article under a publishing agreement with the author(s) or other rightsholder(s); author self-archiving of the accepted manuscript version of this article is solely governed by the terms of such publishing agreement and applicable law.