



Bayesian nonparametric priors for hidden Markov random fields

Hongliang Lu, Jylyan Arbel, Florence Forbes

► To cite this version:

Hongliang Lu, Jylyan Arbel, Florence Forbes. Bayesian nonparametric priors for hidden Markov random fields. *Statistics and Computing*, 2020, 30, pp.1015-1035. 10.1007/s11222-020-09935-9 . hal-02163046v3

HAL Id: hal-02163046

<https://hal.science/hal-02163046v3>

Submitted on 29 Jul 2020

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

Bayesian nonparametric priors for hidden Markov random fields

Hongliang Lü · Julyan Arbel · Florence Forbes

Received: date / Accepted: date

Abstract One of the central issues in statistics and machine learning is how to select an adequate model that can automatically adapt its complexity to the observed data. In the present paper, we focus on the issue of determining the structure of clustered data, both in terms of finding the appropriate number of clusters and of modelling the right dependence structure between the observations. Bayesian nonparametric (BNP) models, which do not impose an upper limit on the number of clusters, are appropriate to avoid the required guess on the number of clusters but have been mainly developed for independent data. In contrast, Markov random fields (MRF) have been extensively used to model dependencies in a tractable manner but usually reduce to finite cluster numbers when clustering tasks are addressed. Our main contribution is to propose a general scheme to design tractable BNP-MRF priors that combine both features: no commitment to an arbitrary number of clusters and a dependence modelling. A key ingredient in this construction is the availability of a stick-breaking representation which has the three-fold advantage to allowing us to extend standard discrete MRFs to infinite state space, to design a tractable estimation algorithm using variational approximation and to derive theoretical properties on the predictive distribution and the number of clusters of the proposed model. This approach is illustrated on a challenging natural image segmentation task for which it shows good performance with respect to the literature.

Keywords Hidden Markov random fields · Bayesian nonparametrics · Variational approximation · Clustering · Image segmentation · Predictive distribution

1 Introduction

Hidden Markov random field (HMRF) models are widely used for clustering data under spatial constraints. Spatial dependencies are encoded by modelling the cluster labels as a discrete state Markov random field (MRF) such as Ising (two clusters or states) or Potts (more than two clusters) model [11,35]. HMRF can be seen as spatial extensions of independent mixture models. As for standard mixtures, one concern is the automatic selection of the proper number of clusters in the data, or equivalently the number of states in the HMRF. In the independent data case, several criteria exist to select this number automatically based on penalized likelihoods (*e.g.*, AIC, BIC, ICL, etc.) and have been extended in the HMRF framework using variational approximation [18]. They require running several models with different cluster numbers so as to choose the best one, with a potential waste of computational effort as all the other models are usually discarded. Other techniques use a fully Bayesian setting including a prior on the number of components. The most celebrated method in this case is reversible jump Markov chain Monte Carlo [20]. Although simplifications in the inference have been proposed recently in [24], the computational cost of reversible jump techniques remains considerably high.

In the present work, we investigate alternatives based on Bayesian nonparametric (BNP) methods. In particular, Dirichlet process mixture (DPM) models have emerged as promising candidates for clustering applications where the number of clusters is unknown. Nevertheless, applications of DPMs involve observations which are assumed to be independent. For more complex tasks such as unsupervised image segmentation with spatial relationships or dependencies between the observations, DPMs are not satisfactory. Therefore, we propose to

introduce MRF dependencies between data points in BNP models, and we term the resulting model BNP-MRF. This requires to extend finite state space MRF models to an infinite number of states. We show that this can be achieved by incorporating a stick-breaking scheme in an MRF formulation more general than the standard Potts model commonly used.

The addition of MRF dependencies between data points in BNP models raises the question of how they impact the natural clustering and rich-get-richer properties of BNP priors. We answer this question by providing theoretical results about two quantities of interest for BNP priors: the predictive distribution, that represents the distribution of one datum conditional on previous observations, and the number of clusters induced by a BNP-MRF prior.

For a practical illustration of the BNP-MRF combination, we detail the example of the Pitman-Yor process prior for which we provide a variational Bayes Expectation-Maximization (VBEM) algorithm. Although the model and theoretical results are more general, the derivation of an explicit variational inference requires specifications.

The links to other similar attempts is reviewed in Section 2. The proposed BNP-MRF model is explained in Section 3 and theoretical properties are investigated in Section 4. The model implementation using variational approximation is detailed in Section 5. An illustration of its performance on an image segmentation task on simulated and real data is provided in Section 6. A conclusion and perspectives are provided in Section 7, while an appendix containing proofs and additional computational derivations ends the paper.

2 Related work

Attempts to build countably infinite state space MRF models using BNP priors have already appeared in the literature. In particular, we can distinguish attempts such as [12, 13, 33] from the work in [2, 27, 42, 34]. The approach in [12, 13, 33] differs in that it is not based on a generalization of the Potts model but on a transformation of an inference algorithm. More specifically in [12, 13], a standard mean field approximation is first considered and then transformed to account for an infinite number of states. In that sense it is closer to an Iterated Conditional Mode (ICM) algorithm [7], but does not provide a spatial generalization of DPMs. Typically, the simple Potts model considered in [12, 13] cannot be extended to an infinite number of states as it will become clear in our Section 3.2. Other attempts include the work in [22], but there the number of states is known to be three and the Dirichlet process (DP) prior is used instead to model intensity distributions non-parametrically. Segmentation with spatially dependent Pitman–Yor processes (PY) has also been considered in [36], but using Gaussian processes.

We build on the approach in [2] which differs from [27, 42, 34] which all use a partition model representation. In particular, [42] generalizes [27] and proposes a more efficient Markov chain Monte Carlo (MCMC) inference by means of the Swendsen–Wang algorithm, while [34] extends this idea to hierarchical DP priors for multiple image segmentation. In contrast to [27, 42, 34], we propose to use a stick-breaking-based scheme for the mixing weights, thus providing a more comprehensive representation than partition models which integrate out the process. In addition, stick-breaking representations lead naturally to variational approximations for performing inference [8]. The advantage is to reduce the computational cost in complex data clustering without suffering from label switching complications. In other words, in our approach the MRF is imposed internally in the BNP mechanics leading to well defined infinite state HMRF models. This construction is valid for any stick-breaking representation. We show how it can be implemented for the DP and PY priors, and provide references for extensions to larger classes of BNP priors.

3 BNP-MRF mixture models

The clustering task is addressed through a missing data model that includes a set $\mathbf{y} = (y_1, \dots, y_n)$ of observed variables from $\mathbb{R}^d$ and a set $\mathbf{z} = (z_1, \dots, z_n)$ of missing (also called hidden) variables whose joint distribution $p(\mathbf{y}, \mathbf{z} \mid \boldsymbol{\Theta})$ is governed by a set of parameters denoted by $\boldsymbol{\Theta}$ and possibly by additional hyperparameters ϕ not specified in the notation. The latter ones are usually fixed and not considered at first. Typically, the z_i 's corresponding to group memberships (or labels), take their values in $\{1, \dots, K\}$ where K is the number of clusters or groups. We shall denote by $\mathcal{Z} = \{1, \dots, K\}^n$ the set in which $\mathbf{z}$ takes its values and by $\underline{\Theta}$ the parameter space. To account for dependencies between the z_i 's, $\mathbf{z}$ can be modeled as a discrete MRF. If in addition, the y_j 's are independent conditionally on $\mathbf{z}$, the joint distribution $p(\mathbf{y}, \mathbf{z} \mid \boldsymbol{\Theta})$ is referred to as an HMRF model. In this case, the conditional distribution $p(\mathbf{z} \mid \mathbf{y}, \boldsymbol{\Theta})$ is also an MRF. For clustering dependent data into K groups, the most commonly used MRF is the so-called Potts model [11, 35].

As already mentioned, our goal is to bypass the issue of selecting the number K of clusters by considering a countably infinite number of them while allowing MRF dependencies between the $\mathbf{y}_i$'s. The construction of the proposed model is explained starting from the link between standard finite mixtures and Dirichlet process mixtures. Basic DP principles and notations are recalled in Section 3.1. The extension of finite state space MRF to a countably infinite number of states is given in Section 3.2 and the resulting BNP-MRF mixture models is summarized in Section 3.3.

3.1 From finite mixtures to DP mixture models

A generative approach to clustering consists of picking one of K clusters from a multinomial distribution with weights parameter $\boldsymbol{\pi} = (\pi_1, \dots, \pi_K)$ and then to generate a data point y from a cluster specific distribution $p(y | \theta_k^*)$ with cluster specific parameter θ_k^* . This yields a finite mixture model

$$p(y | \boldsymbol{\theta}^*, \boldsymbol{\pi}) = \sum_{k=1}^K \pi_k p(y | \theta_k^*) \quad (1)$$

where $\boldsymbol{\theta}^* = (\theta_1^*, \dots, \theta_K^*)$ and $\boldsymbol{\pi}$ are the parameters. For instance, for Gaussian mixtures, $\theta_k^* = (\mu_k, \Sigma_k)$ and $p(y | \theta_k^*)$ is a Gaussian distribution with mean μ_k and covariance matrix Σ_k , denoted by $\mathcal{N}(\mu_k, \Sigma_k)$ or $\mathcal{N}(y | \mu_k, \Sigma_k)$ when referring to the probability density function (pdf). The observations $(y_1, \dots, y_n)$ are therefore *i.i.d.* and generated from the same mixture (1). It follows that the k th cluster is by definition the set of data points arising from the k th mixture component. This is usually expressed by introducing for each y_j an additional hidden variable Z_j that takes its values in $\{1, \dots, K\}$, so that $p(z_j = k | \boldsymbol{\pi}) = \pi_k$. Another way to obtain a sample from a finite mixture model consists of defining a discrete measure $G = \sum_{k=1}^K \pi_k \delta_{\theta_k^*}$ and then of considering the following hierarchical representation, for all $j = 1, \dots, n$,

$$\begin{aligned} \theta_j &| G \stackrel{\text{iid}}{\sim} G, \\ y_j &| \theta_j \stackrel{\text{ind}}{\sim} p(\cdot | \theta_j). \end{aligned}$$

The subset of θ_j 's that are equal to θ_k^* corresponds to the y_j 's in the k th cluster.

In a Bayesian setting, in addition, a prior distribution is placed on $\boldsymbol{\theta}^*$ and $\boldsymbol{\pi}$. The most common choice for $\boldsymbol{\pi}$ is the Dirichlet distribution $\text{Dir}(\alpha_1, \dots, \alpha_K)$ depending on a vector of positive parameters $\boldsymbol{\alpha} = (\alpha_1, \dots, \alpha_K)$. The choice of the prior on $\boldsymbol{\theta}^*$ (denoted by G_0) is model-specific, usually following a conjugate prior such as a Normal inverse-Wishart distribution for Gaussian mixture models. Other cases are possible and tractable (*e.g.* [14]). It follows the hierarchical representation:

$$\theta_1^*, \dots, \theta_K^* | G_0 \sim G_0, \quad (2)$$

$$\boldsymbol{\pi} | \boldsymbol{\alpha} \sim \text{Dir}(\alpha_1, \dots, \alpha_K), \quad (3)$$

$$G = \sum_{k=1}^K \pi_k \delta_{\theta_k^*}, \quad (4)$$

$$\theta_j | G \stackrel{\text{iid}}{\sim} G, j = 1, \dots, n, \quad (5)$$

$$y_j | \theta_j \stackrel{\text{ind}}{\sim} p(\cdot | \theta_j) j = 1, \dots, n.$$

To become non-parametric, a first approach is to consider an infinite number of π_k 's. Using an infinite number of random variables $\boldsymbol{\tau} = (\tau_1, \tau_2, \dots)$ on $[0, 1]$, we can construct an infinite number of π_k 's that sum to one as follows:

$$\pi_1(\boldsymbol{\tau}) = \tau_1 \quad \text{and} \quad \pi_k(\boldsymbol{\tau}) = \tau_k \prod_{l=1}^{k-1} (1 - \tau_l), \quad k = 2, 3, \dots$$

A necessary and sufficient condition to guarantee that these π_k 's sum to 1 almost surely is that the expectation $\mathbb{E}[\prod_{l=1}^k (1 - \tau_l)]$ tends to 0 as k tends to ∞ . In particular, if $\tau_1, \tau_2, \dots$ are *i.i.d.* it suffices that $p(\tau_1 > 0) > 0$. The proof and additional details can be found in Lemma 3.4 of [19]. The intuition behind this construction, referred to as *stick-breaking* and proposed by [31], is that it consists of recursively breaking a unit-length stick

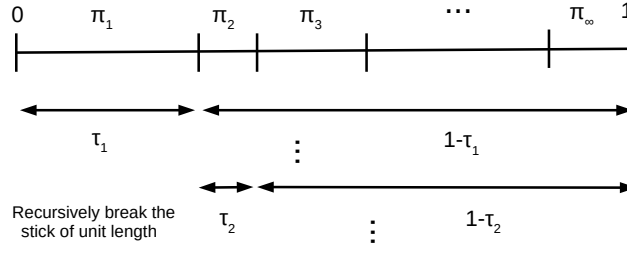


Fig. 1 Illustration of the stick-breaking representation.

as shown in Fig. 1. It follows an explicit formula for the π_k 's. Hence, the τ_k 's simulation replaces step (3), and G in (4) can be replaced by

$$G = \sum_{k=1}^{\infty} \pi_k(\tau) \delta_{\theta_k^*}.$$

We can also add after step (5) the fact that $z_j = k$ if $\theta_j = \theta_k^*$ and replace the last step by $y_j | z_j, \theta^* \stackrel{\text{ind}}{\sim} p(\cdot | \theta_{z_j}^*)$. Then the distributions of the τ_k 's need to be specified. The Dirichlet process [17], denoted by $\text{DP}(G_0, \alpha)$, is characterized by a base distribution G_0 and a positive scaling parameter α . Its stick-breaking representation corresponds to *i.i.d* τ_k 's that follow the same beta $\mathcal{B}(1, \alpha)$ distribution [21]. All together, using the same notation G_0 for the prior of each θ_k^* simulated as *i.i.d* variables, it comes the following hierarchical representation:

$$\theta_k^* | G_0 \stackrel{\text{iid}}{\sim} G_0, \quad k = 1, 2, \dots, \quad (6)$$

$$\tau_k | \alpha \stackrel{\text{iid}}{\sim} \mathcal{B}(1, \alpha), \quad k = 1, 2, \dots,$$

$$\pi_k(\tau) = \tau_k \prod_{l=1}^{k-1} (1 - \tau_l), \quad k = 1, 2, \dots, \quad (7)$$

$$G = \sum_{k=1}^{\infty} \pi_k(\tau) \delta_{\theta_k^*}, \quad (8)$$

$$\theta_j | G \stackrel{\text{iid}}{\sim} G, \text{ and } z_j = k \text{ if } \theta_j = \theta_k^* \quad (9)$$

$$y_j | z_j, \theta^* \stackrel{\text{ind}}{\sim} p(\cdot | \theta_{z_j}^*). \quad (10)$$

The above hierarchical representation corresponds to a countably infinite mixture model referred to as a Dirichlet process mixture (DPM) model. It is an explicit characterization of the DP (Eq. (6) to (8)) and of the DPM (Eq. (6) to Eq. (10)) using a stick-breaking construction. The stick-breaking representation will be particularly useful in our study for both the definition of our model (Sections 3.2 and 3.3) and its estimation (Section 5).

3.2 Infinite MRF priors

The explicit use of the labels $\mathbf{z} = (z_1, \dots, z_n)$ in the DPM construction above makes it closer to clustering generative models and opens the way to an HMRF extension. Such a generalization is only possible from Potts models with an external field parameter. In the finite state space case, an MRF model is defined using a dependence structure coded via a graph $\mathcal{G}$ whose nodes correspond to the variables. A K -state Potts model with an external field, defined over $\mathbf{z} = (z_1, \dots, z_n)$ with for all $j = 1, \dots, n$, $z_j \in \{1, \dots, K\}$, corresponds to the following pdf,

$$p(\mathbf{z}; \beta, \mathbf{v}) \propto \exp \left(\sum_{j=1}^n v_{z_j} + \beta \sum_{i \sim j} \delta_{(z_i = z_j)} \right), \quad (11)$$

where $i \sim j$ means that i and j are neighbors, *i.e.* linked by an edge, in the considered dependence structure described by graph $\mathcal{G}$, $\delta_{(z_i = z_j)}$ is the indicator function which is 1 if $z_i = z_j$ and 0 otherwise, β is a positive scalar interaction parameter and $\mathbf{v} = (v_1, \dots, v_K)$ represents an additional external field parameter where each v_k is a scalar. The distribution (11) is insensitive to an addition of the same constant to all the v_k 's. Such

non-identifiability can be overcome by an additional constraint on $\mathbf{v}$ such as requiring $\sum_{k=1}^K \pi_k = 1$ with $v_k = \log \pi_k$. The Potts model in (11) can then be rewritten as

$$p(\mathbf{z}; \beta, \boldsymbol{\pi}) \propto \left(\prod_{j=1}^n \pi_{z_j} \right) \exp \left(\beta \sum_{i \sim j} \delta_{(z_i = z_j)} \right). \quad (12)$$

In the finite state space case, we can equivalently use the Gibbs representation,

$$p(\mathbf{z}; \beta, \boldsymbol{\pi}) \propto e^{V(\mathbf{z}; \beta, \boldsymbol{\pi})}, \quad (13)$$

where $V(\mathbf{z}; \beta, \boldsymbol{\pi}) := \sum_{j=1}^n \log \pi_{z_j} + \beta \sum_{i \sim j} \delta_{(z_i = z_j)}$ is often referred to as the *energy function*. The first sum in V represents the first order potentials while the second sum represents the second order potentials. In the finite state space case, the Hammersley–Clifford theorem [7] applied to the Gibbs representation (13) entails that the distribution in (11) is a Markov random field. What is interesting about formulas (11) and (12) is that they do not involve the number of states K . As long as a stick-breaking construction is available, we can consider a countably infinite number of probabilities π_k that sum to one, *i.e.*, $\sum_{k=1}^{\infty} \pi_k = 1$ and define the same energy function V as before but over an infinite countable set of states. Using the Gibbs representation (13), the Hammersley–Clifford theorem still holds if we can show that $\sum_{\mathbf{z}} e^{V(\mathbf{z}; \boldsymbol{\pi}, \beta)} < \infty$, where the sum runs over all n -uples of positive integers $\mathbf{z} \in \{1, 2, \dots\}^n$. Note that this latter condition that is automatically satisfied in the finite state space case (for reasonable potential functions), may not be satisfied in the infinite case. However, the stick-breaking representation of $\boldsymbol{\pi}$ ensures this property since:

$$\sum_{\mathbf{z}} e^{V(\mathbf{z}; \beta, \boldsymbol{\pi})} \stackrel{(a)}{\leq} \left(\sum_{\mathbf{z}} \prod_{j=1}^n \pi_{z_j} \right) e^{\beta \frac{n(n-1)}{2}} \stackrel{(b)}{=} e^{\beta \frac{n(n-1)}{2}} < \infty$$

where we used for (a) the fact that $n(n-1)/2$ is the maximum number of neighbors among n observations (complete dependence or graph), while (b) comes from $\sum_{\mathbf{z}} \prod_{j=1}^n \pi_{z_j} = \left(\sum_{k=1}^{\infty} \pi_k \right)^n = 1$. It follows that $p(\mathbf{z}; \beta, \boldsymbol{\pi})$, in the infinite state space case, is still a valid probability distribution and is an MRF by the Hammersley–Clifford theorem. Such a generalization is possible because of the presence of the external field parameters π_k that satisfy $\sum_{k=1}^{\infty} \pi_k = 1$ as ensured by the stick-breaking construction. A standard Potts model with equal or no external field parameters cannot be as simply extended to an infinite countable state space because in the K -state case this Potts model is equivalent to $\pi_k = 1/K$ for all k which possesses a degenerate limit when K tends to infinity.

3.3 BNP-MRF mixture models

The stick-breaking representation amounts to identifying a set of random variables $\boldsymbol{\tau} = (\tau_k)_{k=1}^{\infty}$ with each $\tau_k \in [0, 1]$ and so that the weights π_k are defined by (7). Then the Potts model construction (12) is valid for any set of parameters $\boldsymbol{\tau} = (\tau_k)_{k=1}^{\infty}$ with each $\tau_k \in [0, 1]$. Bayesian non-parametric priors specify a prior distribution on τ_k 's. For instance, as already mentioned for the DP stick-breaking, all τ_k 's are independent and identically distributed according to a $\mathcal{B}(1, \alpha)$ distribution. For the Pitman–Yor (PY) process [29], the τ_k 's are independent but not identically distributed with

$$\tau_k \mid \alpha, \sigma \stackrel{\text{ind}}{\sim} \mathcal{B}(1 - \sigma, \alpha + k\sigma) \quad \text{for } k = 1, 2, \dots, \quad (14)$$

where $\sigma \in (0, 1)$ is a discount parameter and α a concentration parameter $\alpha > -\sigma$. The PY is a two-parameter generalisation of the DP which allows to control the tail behavior when modeling data with either exponential or power-law tails [21, 29]. When $\sigma = 0$, the PY reduces to a DP. More general stick-breaking representations are possible (*e.g.*, for Gibbs-type priors [15, 19] or homogeneous normalized random measures with independent increments (NRMIs) [16]) but the Pitman–Yor case provides a clear interpretation in terms of number of clusters. The rich-gets-richer property of the DP is preserved meaning that there are a small number of large clusters, but there is also a large number of small clusters with parameter σ decreasing the probability that observations join small clusters. The PY yields a power-law behavior which can make it more suitable for a number of applications. In other words, the number of clusters grows as $\mathcal{O}(n^\sigma)$ for the PY while it grows more slowly at $\mathcal{O}(\log n)$ for the DP.

The extension we propose is therefore to augment the original HMRF formulation with additional variables $(\tau_k)_{k=1}^\infty$. We refer to it as the BNP-MRF mixture model. It corresponds to the following hierarchical construction written here in the PY case:

$$\theta_k^* \mid G_0 \stackrel{\text{iid}}{\sim} G_0, \quad k = 1, 2, \dots, \quad (15)$$

$$\tau_k \mid \alpha, \sigma \stackrel{\text{iid}}{\sim} \mathcal{B}(1 - \sigma, \alpha + k\sigma), \quad k = 1, 2, \dots, \quad (16)$$

$$\pi_k(\boldsymbol{\tau}) = \tau_k \prod_{l=1}^{k-1} (1 - \tau_l), \quad (17)$$

$$p(\mathbf{z} \mid \boldsymbol{\tau}; \beta) \propto \left(\prod_{j=1}^n \pi_{z_j}(\boldsymbol{\tau}) \right) \exp \left(\beta \sum_{i \sim j} \delta_{(z_i = z_j)} \right), \quad (18)$$

$$y_j \mid z_j, \boldsymbol{\theta}^* \stackrel{\text{iid}}{\sim} p(y_j \mid \theta_{z_j}^*). \quad (19)$$

The prior on τ_k 's from (16) can be adapted to more general classes of BNP priors, see for example Theorem 14.23 of [19] for Gibbs-type priors, and [16] for NRMIs. Importantly, in the BNP-MRF model above, the θ_j 's and z_j 's are not *i.i.d* conditionally on G anymore. The joint distribution (18) on $\mathbf{z} = (z_1, \dots, z_n)$ induces a joint distribution on $(\theta_1, \dots, \theta_n)$ using that $\theta_j = \theta_{z_j}^*$. If we still denote for simplicity by G this joint distribution, we can define it in a similar manner as in the *i.i.d.* case, using its conditional specifications,

$$\theta_j \mid \theta_{\mathcal{N}_j}; G \sim \sum_{k=1}^{\infty} p(z_j = k \mid z_{\mathcal{N}_j}, \boldsymbol{\tau}; \beta) \delta_{\theta_k^*},$$

where $\mathcal{N}_j$ denotes the neighbors of j in the graph dependence structure $\mathcal{G}$ and the $p(z_j = k \mid z_{\mathcal{N}_j}, \boldsymbol{\tau}; \beta)$'s are the conditional specifications of (18).

In Section 5.2, we detail the case when cluster specific distributions are Gaussian, with $\theta_k^* = (\mu_k, \Sigma_k)$ and $p(y_j \mid \theta_k^*) = \mathcal{N}(y_j \mid \mu_k, \Sigma_k)$.

4 Predictive distribution and probabilistic properties of the BNP-MRF prior

In this section, we provide theoretical results about two quantities of interest for Bayesian nonparametric priors: the predictive distribution, that represents the distribution of one datum conditional on previous observations, and the number of clusters induced by a BNP-MRF prior. We then turn to an empirical, simulation-based investigation of the distribution of the cluster sizes, for which we do not have theoretical results.

4.1 Predictive distribution

We consider data of varying sample size, and denote by $\mathcal{G}_n$ the subgraph of $\mathcal{G}$ induced by node $\{1, \dots, n\}$. We focus on the large class of Gibbs-type priors [15], of which the DP and PY are special cases. Consider n observations $(\theta_1, \dots, \theta_n)$ sampled from a BNP-MRF prior Eq (15)-(18) but using a Gibbs-type prior instead of PY prior (16). We are interested in the predictive distribution of observation θ_{n+1} conditional on $(\theta_1, \dots, \theta_n)$, but unconditional on G . With a BNP-MRF prior, this predictive distribution depends on the structure of the graph $\mathcal{G}$, more specifically on the neighbors of θ_{n+1} . Denote by K_n the number of clusters in $(\theta_1, \dots, \theta_n)$, by $(\theta_1^*, \dots, \theta_{K_n}^*)$ their K_n different values¹ and by $(n_1, \dots, n_{K_n})$ their size. We first consider the Gibbs-type prior case without the addition of a Markov component. The predictive distribution [19] is given by,

$$p(\theta_{n+1} \mid \theta_1, \dots, \theta_n) = \frac{V_{n+1, K_n+1}}{V_{n, K_n}} G_0 + \frac{V_{n+1, K_n}}{V_{n, K_n}} \sum_{\ell=1}^{K_n} (n_\ell - \sigma) \delta_{\theta_\ell^*} \quad (20)$$

where the triangular array of non-negative parameters $V_{n,k}$, $1 \leq k \leq n$, satisfy the backward recurrence relation

$$V_{n,k} = (n - \sigma k) V_{n+1, k} + V_{n+1, k+1}, \quad (21)$$

¹ Note that the notation introduced for the different θ_j^* differs from that devoted to the stick-breaking variables, θ_j^* .

with $V_{1,1} = 1$. This predictive can be specialized to the PY case with

$$V_{n,k} = \frac{\sigma^k (1 + \frac{\alpha}{\sigma})_{(k-1)}}{(1 + \alpha)_{(n-1)}},$$

where $(a)_{(x)} := \Gamma(a + x)/\Gamma(a)$ denotes the rising factorial. It follows

$$p(\theta_{n+1} \mid \theta_1, \dots, \theta_n) = \frac{\alpha + \sigma K_n}{\alpha + n} G_0 + \frac{1}{\alpha + n} \sum_{\ell=1}^{K_n} (n_\ell - \sigma) \delta_{\theta_\ell^*}, \quad (22)$$

while the case of the DP is obtained by setting $\sigma = 0$ above.

For the sake of simplicity, we propose to use the labels notation $z_{1:n} = (z_1, \dots, z_n)$ defined so that $z_j = \ell$ when $\theta_j = \theta_\ell^*$, we denote by $\{z_1, \dots, z_n\}$ the set of label values which includes only K_n different labels. In the Gibbs-type prior case, it is clear from (20) that

$$\begin{aligned} p(z_{n+1} \mid z_{1:n}) &= \frac{V_{n+1, K_n+1}}{V_{n, K_n}} \quad \text{if } z_{n+1} \notin \{z_1, \dots, z_n\}, \\ p(z_{n+1} = \ell \mid z_{1:n}) &= \frac{V_{n+1, K_n}}{V_{n, K_n}} (n_\ell - \sigma) \quad \text{if } \ell \in \{z_1, \dots, z_n\}. \end{aligned} \quad (23)$$

The next proposition indicates how the predictive is impacted by the addition of a Markov dependence. The neighbors of θ_{n+1} in $\mathcal{G}_n$ is denoted by $\mathcal{N}_{n+1}$ and $\tilde{n}_\ell$ is the number of neighbors of θ_{n+1} which belong to cluster ℓ , hence satisfying $\tilde{n}_\ell \leq n_\ell$. Also $z_{\mathcal{N}_{n+1}} = \{z_i, i \in \mathcal{N}_{n+1}\}$ denotes the labels in the neighborhood. The proof of the proposition is given in Appendix A.1.

Proposition 1 (Predictive distribution of a Gibbs-MRF prior) *The predictive distribution for a Gibbs-MRF prior is given by*

$$p(\theta_{n+1} \mid \theta_1, \dots, \theta_n) = \frac{V_{n+1, K_n+1}}{V_{n, K_n} + V_{n+1, K_n} \boldsymbol{\eta}_{n+1}} G_0 + \frac{V_{n+1, K_n}}{V_{n, K_n} + V_{n+1, K_n} \boldsymbol{\eta}_{n+1}} \sum_{\ell=1}^{K_n} \boldsymbol{\lambda}_{n+1, \ell} \delta_{\theta_\ell^*} \quad (24)$$

where

$$\begin{aligned} \boldsymbol{\eta}_{n+1} &= \boldsymbol{\eta}_{n+1}(\sigma, \beta) = \sum_{\ell \in z_{\mathcal{N}_{n+1}}} (n_\ell - \sigma) (e^{\beta \tilde{n}_\ell} - 1), \\ \boldsymbol{\lambda}_{n+1, \ell} &= \boldsymbol{\lambda}_{n+1, \ell}(\sigma, \beta) = (n_\ell - \sigma) e^{\beta \tilde{n}_\ell \delta_{\mathcal{N}_{n+1}}(\ell)}. \end{aligned}$$

and $\delta_{\mathcal{N}_{n+1}}(\ell)$ is 1 when ℓ is a label present in the neighborhood of θ_{n+1} and 0 otherwise.

Remark 1 When $\beta = 0$, $\boldsymbol{\eta}_{n+1}(\sigma, 0) = 0$ and $\boldsymbol{\lambda}_{n+1, \ell}(\sigma, 0) = n_\ell - \sigma$ so that the Gibbs-type prior predictive (20) is recovered. In contrast, for $\beta > 0$, the above predictive specialized to the PY-MRF case is,

$$p(\theta_{n+1} \mid \theta_1, \dots, \theta_n) = \frac{\alpha + \sigma K_n}{\alpha + n + \boldsymbol{\eta}_{n+1}} G_0 + \frac{1}{\alpha + n + \boldsymbol{\eta}_{n+1}} \sum_{\ell=1}^{K_n} \boldsymbol{\lambda}_{n+1, \ell} \delta_{\theta_\ell^*}, \quad (25)$$

while the case of the DP-MRF is obtained by setting $\sigma = 0$. Comparing the probability of a new draw for a Gibbs-type prior, $\frac{V_{n+1, K_n+1}}{V_{n, K_n}}$, with that of a new draw for a Gibbs-MRF prior, $\frac{V_{n+1, K_n+1}}{V_{n, K_n} + V_{n+1, K_n} \boldsymbol{\eta}_{n+1}}$, we see that the MRF has the effect of reducing this probability. In the PY case, this increase corresponds to increasing the sample size from n to $n + \boldsymbol{\eta}_{n+1}$ when $\beta > 0$, where $\boldsymbol{\eta}_{n+1}$ can be quite a large number. More specifically for a label ℓ in the neighborhood of z_{n+1} , the weight of each previous observations with label ℓ (in the neighborhood or not) is multiplied by a factor $(e^{\beta \tilde{n}_\ell} - 1)$. The effect is then all the more important as β is large and as n_ℓ is large.

4.2 Number of clusters

The predictive distribution (24) provides in turn the following lower bounds on the prior expectation of the number of clusters. The proof of Proposition 2 is given in Appendix A.2.

Proposition 2 (Lower bound for expected number of clusters) *Assume that the graph $\mathcal{G}$ has maximal degree D . Then the expected prior number of clusters for a BNP-MRF distribution has the following lower bound*

$$\mathbb{E}[K_n] \gtrsim \frac{\alpha}{e^{D\beta}} \log n$$

for the Dirichlet process and

$$\mathbb{E}[K_n] \gtrsim cn^{\sigma e^{-D\beta}}, \quad (26)$$

for the Pitman–Yor process, with some positive constant c , and where $a_n \gtrsim b_n$ stands for $\limsup a_n/b_n \geq 1$.

Remark 2 We do not have a proof for the general case of Gibbs-type priors, but we conjecture that the same power-law lower bound (26) as for PY holds.

Note that the MRF component of a BNP prior can only *reduce* the prior expected number of clusters. For instance, for the DP with a simple graph where the first two nodes are connected, we have

$$\mathbb{E}[K_2] = 1 + \frac{\alpha}{\alpha + e^\beta} \leq 1 + \frac{\alpha}{\alpha + 1} = \mathbb{E}[K_2; \beta = 0]$$

where the last two terms above correspond the expectation of K_2 for a DP, *i.e.* when $\beta = 0$. Thus natural upper bounds that complement the lower bounds of Proposition 2 are given by

$$\mathbb{E}[K_n] \lesssim \alpha \log n$$

for the Dirichlet process and

$$\mathbb{E}[K_n] \lesssim \frac{\Gamma(\alpha + 1)}{\sigma \Gamma(\alpha + \sigma)} n^\sigma$$

for the Pitman–Yor process (see [28]).

4.3 Clusters sizes

We conclude this section by empirically investigating the cluster sizes. To this purpose, let us introduce the notation $K_{j,n}$, $j = 1, \dots, n$, for the number of clusters of size j in a sample of n from the BNP-MRF prior. Note that $\sum_{j=1}^n K_{j,n} = K_n$ and $\sum_{j=1}^n j K_{j,n} = n$.

The distribution of cluster sizes is mentioned for instance in [40] in the context of Dirichlet process and Pitman–Yor priors (without MRF component), with their asymptotics when n tend to infinity given by their Equations (6) for the DP and (8) for PY. Those equations are best illustrated in a log-log scale as they yield, for n large enough:

$$\log(\mathbb{E}[K_{j,n}]) \approx c - (1 + \sigma) \log j, \quad (27)$$

where c is a constant. This is easily seen for the DP from Equation (6) of [40]. For the PY, this comes from the approximation $\frac{\Gamma(j - (1 + \sigma))}{\Gamma(j)} \approx j^{-(1 + \sigma)}$ for large j in Equation (8) of [40]. Hence in a log-log scale plot, the cluster size distribution for large n in the DP case has a slope equal to -1 , while a PY features a steeper slope of $-(1 + \sigma)$. Note that [40] also illustrate (see their Figure 2) that the asymptotic behaviours are close to the finite sample behaviours for small j but require n to be greater than 10^5 for larger j values. For $n = 1000$ the asymptotic formula do not provide a so close fit to the true distributions.

In order to empirically assess the cluster sizes for the DP and PY specifications and their MRF counterparts, we have sampled 10^5 images of size $32 \times 32 = 1024$ from the predictive distribution (20) of Proposition 1. The number of neighbors is set to 4 and various values of β and σ are investigated while α is set to 2. The approximate expectations $\mathbb{E}[K_{j,n}]$ for varying j size can be seen in Figure 2 where the Monte Carlo means have been further smoothed over j for visualization. The effect of β can be clearly visualized. Larger β values tend to create larger clusters with maximum sizes while the number of clusters of intermediate sizes tend to decrease. An interpretation is that for large enough β , above a certain size, clusters tend to attract all samples so that eventually they reach a close to maximum size. Also it appears that β has some effect when combined with σ on

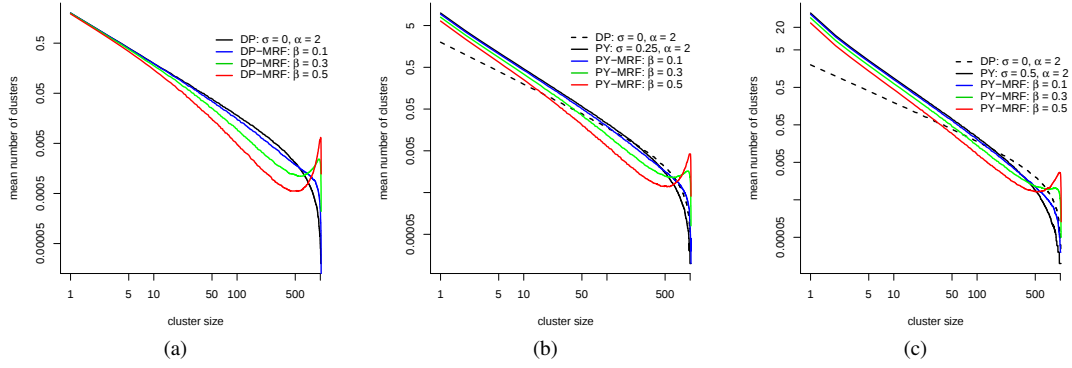


Fig. 2 Empirical cluster sizes obtained over 10^5 images of size $32 \times 32 = 1024$ with 4 neighbors, simulated using the BNP-MRF predictive distribution (20) with Pitman-Yor specifications with $\alpha = 2$, $\sigma \in \{0, 0.25, 0.5\}$ and $\beta \in \{0, 0.1, 0.3, 0.5\}$. The curves correspond to the Monte Carlo approximations of the expected number of clusters of size j with an additional smoothing over j for visualization.

the intercept of the slopes, see Figure 2 (b,c). As regards the final peaks observed when β is greater than 0.1, we believe it is not an artefact. It may be the result of a competition in the BNP-MRF model between two opposite effects. While in pure BNP models ($\beta = 0$) the probability of a large cluster decreases to 0 with its size (see the black curves in Figure 2), configurations with larger clusters have higher probability in a Potts model and all the higher as β increases. In other words, for small j values, the high number of configurations with $K_{n,j}$ clusters of size j compensates the low probability of each of these configurations. While for larger j values, the number of such configurations decreases but their probability is much higher. This compensation based on the fact that a Potts model tends to favor large clusters, depends on the interaction parameter β and is only effective and visible for large enough β (here above 0.1) and large enough clusters. Empirically, the peaks in the BNP-MRF curves correspond to the drop in the gradients of the pure BNP curves, that is when the expected number of clusters of size j drops significantly faster.

5 Inference using variational approximation

Sampling based inference (MCMC) for a similar BNP-MRF model has been proposed in [27,42] for the case of a DP prior. As an alternative, we propose a variational approximation that is facilitated by the stick-breaking representation. For that purpose, we shall briefly recall the variational principle.

5.1 Variational Bayesian Expectation Maximization

The clustering task consists primarily of estimating the unknown labels $\mathbf{z} = (z_1, \dots, z_n)$ from observed $\mathbf{y} = (y_1, \dots, y_n)$ assuming a joint distribution $p(\mathbf{y}, \mathbf{z} \mid \boldsymbol{\theta}; \phi)$ governed by a set of parameters denoted by $\boldsymbol{\theta}$ and often by additional hyperparameters ϕ . However to perform good label estimation, the parameters $\boldsymbol{\theta}$ values (and hyperparameters ϕ) have to be available. A natural approach for parameter estimation is based on maximum likelihood, where $\boldsymbol{\theta}$ is estimated by $\hat{\boldsymbol{\theta}} = \arg \max_{\boldsymbol{\theta} \in \Theta} p(\mathbf{y} \mid \boldsymbol{\theta})$. Then an estimate of $\mathbf{z}$ can be obtained by maximizing $p(\mathbf{z} \mid \mathbf{y}, \hat{\boldsymbol{\theta}})$. However, $p(\mathbf{y} \mid \boldsymbol{\theta})$ is a marginal distribution over the unknown $\mathbf{z}$ variables, so that direct maximum likelihood is intractable in general. The Expectation-Maximization (EM) algorithm [23] is a general iterative technique for maximum likelihood estimation in the presence of unobserved latent variables or missing data. An EM iteration consists of two steps usually referred to as the E-step in which the expectation of the so-called complete log-likelihood is computed and the M-step in which this expectation is maximized over $\boldsymbol{\theta}$. An equivalent way to define EM is the following. As discussed in [26], EM can be viewed as an alternating maximization procedure of a function $\mathcal{F}_0$ defined, for any probability distribution $q_{\mathbf{Z}}$ on $\mathcal{Z}$ by

$$\begin{aligned} \mathcal{F}_0(q_{\mathbf{Z}}, \boldsymbol{\theta}, \phi) &= \sum_{\mathbf{z} \in \mathcal{Z}} q_{\mathbf{Z}}(\mathbf{z}) \log p(\mathbf{y}, \mathbf{z} \mid \boldsymbol{\theta}; \phi) + I[q_{\mathbf{Z}}] \\ &= \mathbb{E}_{q_{\mathbf{Z}}} \left[\log \frac{p(\mathbf{y}, \mathbf{Z} \mid \boldsymbol{\theta}; \phi)}{q_{\mathbf{Z}}(\mathbf{Z})} \right] \end{aligned} \quad (28)$$

where $I[q_{\mathbf{Z}}] = -\mathbb{E}_{q_{\mathbf{Z}}}[\log q_{\mathbf{Z}}(\mathbf{Z})]$ is the entropy of $q_{\mathbf{Z}}$ ($\mathbb{E}_q$ denotes the expectation with regard to q). The function $\mathcal{F}_0$ depends on observations $\mathbf{y}$ which are fixed throughout, hence are omitted from the notation.

Instead of considering only point estimation of Θ , a fully Bayesian approach can be carried out, for instance when prior knowledge on the parameters Θ is available. In this case, we have to compute

$$p(\mathbf{z} | \mathbf{y}) = \int_{\Theta} p(\mathbf{z} | \mathbf{y}, \Theta) p(\Theta | \mathbf{y}) d\Theta \quad (29)$$

Integrating out Θ in this way requires the computation of $p(\Theta | \mathbf{y})$ which is not usually available in closed-form. As an alternative to costly simulation-based methods (MCMC), an EM-like procedure using variational approximation can provide approximations of the marginal posterior distributions $p(\Theta | \mathbf{y})$ and $p(\mathbf{z} | \mathbf{y})$. This approach is referred to as VBEM for Variational Bayesian EM [5]. Note however, that we slightly changed the names of the different steps compared to [5]. While *our* VB-E-Z step is the VB-E step in [5], *our* VB-E- θ step corresponds to the VB-M step in [5]. Our VB-M step has then no equivalent in the formulation of [5], which does not consider hyperparameters estimation. However this step corresponds to what M. Beal in his thesis [6] referred to as the *hyperparameter optimisation* step. The VB-EM procedure presented below therefore corresponds to the *VB-EM with hyperparameters optimisation* of [6], (see for instance Fig. 2.5 or Algorithm 5.3 therein). Let q_Z and q_{Θ} denote respectively distributions over $\mathbf{Z}$ and Θ that will serve as approximations to the true posteriors. Similarly to standard EM, VBEM is maximizing the following *free energy* function defined for any q_Z and q_{Θ} distributions

$$\mathcal{F}(q_Z, q_{\Theta}, \phi) = \mathbb{E}_{q_Z q_{\Theta}} \left[\log \frac{p(\mathbf{y}, \mathbf{Z}, \Theta; \phi)}{q_Z(\mathbf{z}) q_{\Theta}(\Theta)} \right]$$

alternatively over q_Z, q_{Θ} and ϕ . Adding a prior on Θ is formally the same as adding Θ to the missing variables, while the hyperparameters ϕ play the role of the parameters of interest in maximum likelihood estimation.

The alternate maximization of $\mathcal{F}$ yields the VBEM algorithm that decomposes into three steps. It is easy to show, using the Kullback–Leibler (KL) divergence properties, that the maximization over q_Z and q_{Θ} leads to the following E-steps (see Appendix A of [10]). At the r th iteration, using current values $\phi^{(r-1)}$ and $q_{\Theta}^{(r-1)}$, we get the following updating,

$$\begin{aligned} \text{VB-E-Z: } q_Z^{(r)}(\mathbf{z}) &\propto \exp \mathbb{E}_{q_{\Theta}^{(r-1)}} [\log p(\mathbf{y}, \mathbf{z}, \Theta; \phi^{(r-1)})], \\ \text{VB-E-}\Theta: q_{\Theta}^{(r)}(\Theta) &\propto \exp \mathbb{E}_{q_Z^{(r)}} [\log p(\mathbf{y}, \mathbf{Z}, \Theta; \phi^{(r-1)})], \\ \text{VB-M-}\phi: \phi^{(r)} &= \arg \max_{\phi} \mathbb{E}_{q_Z^{(r)} q_{\Theta}^{(r)}} [\log p(\mathbf{y}, \mathbf{Z}, \Theta; \phi)]. \end{aligned}$$

Also, it is worth noticing that if $\mathbf{Y}$ and $\mathbf{Z}$ are independent of ϕ conditionally on Θ , as this is often the case when ϕ gathers the parameters that describe the prior on Θ , then the VB-M-step simplifies into

$$\phi^{(r)} = \arg \max_{\phi} \mathbb{E}_{q_{\Theta}^{(r)}} [\log p(\Theta; \phi)] = \arg \min_{\phi} \text{KL}(q_{\Theta}^{(r)} \| p(\Theta; \phi)). \quad (30)$$

Then $\phi^{(r)}$ is the value that minimizes the KL distance between the prior $p(\Theta; \phi)$ and the variational posterior $q_{\Theta}^{(r)}(\Theta)$. In the conjugate exponential family case, it is known that both distributions belong to the same family [5]. If this family is identifiable it follows that $\phi^{(r)} = \hat{\phi}^{(r)}$ where $\hat{\phi}^{(r)}$ are the variational parameters defining $q_{\Theta}^{(r)}(\Theta)$. A more detailed example is given in Section 5.2.

In practice, we can decide which parameters are treated as genuine parameters Θ or as hyperparameters ϕ , depending on whether some prior knowledge is available only for a subset of the parameters or whether the model has hyperparameters ϕ for which no prior information is available. Also for complex models, q_{Θ} and q_Z may need to be further restricted to simpler forms, such as factorized forms, in order to ensure tractable VB-E-steps. This is illustrated in the next section for the PY-MRF inference.

5.2 VBEM for a PY-MRF mixture model with Gaussian components

The VBEM steps are described for a PY-MRF mixture model as defined in Eq. (16) to (19), with Gaussian distributed observations $\mathbf{y}$. As hyperparameters α and σ may have a significant effect on the growth of the number of clusters with data sample size, it is possible to specify priors on them. For the DP case obtained with $\sigma = 0$, it is suggested in [8] to use a gamma prior over α with two hyperparameters s_1 and s_2 , *i.e.* $\alpha \sim \mathcal{G}(s_1, s_2)$ where s_1 and s_2 can be estimated or fixed. A natural question that arises is then whether one can

also find a tractable prior for the discount parameter σ . We propose to use the following prior that accounts for the constraints $\sigma \in (0, 1)$ and $\alpha > -\sigma$,

$$p(\alpha, \sigma; s_1, s_2, a) = p(\alpha | \sigma; s_1, s_2) p(\sigma; a) \quad (31)$$

where $p(\alpha | \sigma; s_1, s_2)$ is a shifted gamma distribution $\mathcal{SG}(s_1, s_2, \sigma)$ and $p(\sigma; a)$ is a distribution depending on some parameter a not specified for the moment but that can typically be taken as the uniform distribution on the interval $(0, 1)$. Such a shifted gamma distribution is the distribution of a variable $U - \sigma$ where σ is considered as fixed and U follows a gamma distribution $\mathcal{G}(s_1, s_2)$. The pdf of this shifted gamma distribution is obtained from the standard gamma distribution as $p(\alpha | \sigma; s_1, s_2) = \mathcal{G}(\alpha + \sigma; s_1, s_2)$. It follows that the joint distribution of the observed data $\mathbf{y}$ and all latent variables becomes

$$p(\mathbf{y}, \mathbf{z}, \boldsymbol{\Theta}; \phi) = p(\alpha, \sigma; s_1, s_2, a) \prod_{j=1}^n p(y_j | z_j, \boldsymbol{\theta}^*) p(\mathbf{z} | \boldsymbol{\tau}; \beta) \prod_{k=1}^{\infty} p(\tau_k | \alpha, \sigma) \prod_{k=1}^{\infty} p(\theta_k^*; \rho_k),$$

where the notation $\prod_{k=1}^{\infty}$ is a distributional notation, and in addition to the terms already defined in (16) and (18), we specify the likelihood term (19) as a Gaussian distribution $p(y_j | \theta_{z_j}^*) = \mathcal{N}(y_j | \mu_{z_j}, \Sigma_{z_j})$ and the G_0 prior on cluster specific parameters $\theta_k^* = (\mu_k, \Sigma_k)$ as a Normal-inverse-Wishart distribution parameterized by $\rho_k = (m_k, \lambda_k, \Psi_k, \nu_k)$ with a pdf

$$p(\theta_k^*; \rho_k) = \mathcal{NIW}(\mu_k, \Sigma_k; \rho_k) = \mathcal{N}(\mu_k; m_k, \lambda_k^{-1} \Sigma_k) \mathcal{IW}(\Sigma_k; \Psi_k, \nu_k).$$

In the above notation, we consider as hyperparameters the set $\phi = (s_1, s_2, a, \beta, (\rho_k)_{k=1}^{\infty})$ while $\boldsymbol{\Theta} = (\boldsymbol{\tau}, \alpha, \sigma, \boldsymbol{\theta}^*)$.

In most variational approximations, the posteriors are approximated in a factorized form (mean-field approximation). In particular, the intractable MRF posterior on $\mathbf{z}$ is approximated as $q_{\mathbf{z}}(\mathbf{z})$ that factorizes so as to handle intractability due to spatial dependencies, namely

$$q_{\mathbf{z}}(\mathbf{z}) = \prod_{j=1}^n q_{z_j}(z_j).$$

Then, the infinite state space for each z_i is dealt with by choosing a truncation of the state space to a maximum label K [8]. In practice, this consists of assuming that the variational distributions q_{z_j} , for $j = 1, \dots, n$, satisfy $q_{z_j}(k) = 0$ for $k > K$ and that the variational distribution on $\boldsymbol{\tau}$ also factorizes as $q_{\boldsymbol{\tau}}(\boldsymbol{\tau}) = \prod_{k=1}^{K-1} q_{\tau_k}(\tau_k)$, with the additional condition that $\tau_K = 1$. Thus, the truncated variational posterior of parameters $\boldsymbol{\Theta}$ is given by

$$q_{\boldsymbol{\Theta}}(\boldsymbol{\Theta}) = q_{\alpha, \sigma}(\alpha, \sigma) \prod_{k=1}^{K-1} q_{\tau_k}(\tau_k) \prod_{k=1}^K q_{\theta_k^*}(\theta_k^*). \quad (32)$$

These forms of $q_{\mathbf{z}}$ and $q_{\boldsymbol{\Theta}}$ lead to four VB-E steps and three VB-M steps summarized below with details in the Appendix. Set the initial value of ϕ to $\phi^{(0)}$. Then, repeat iteratively the following steps. The iteration index is omitted in the update formulas for simplicity.

VB-E- $\boldsymbol{\tau}$ step

The VB-E- $\boldsymbol{\tau}$ step corresponds to a variational approximation in the exponential family case and results in a posterior from the same family as the prior. It comes for $k = 1, \dots, K$,

$$q_{\tau_k}(\tau_k) = \mathcal{B}(\tau_k; \hat{\gamma}_{k,1}, \hat{\gamma}_{k,2}) \quad (33)$$

with

$$\hat{\gamma}_{k,1} = 1 - \mathbb{E}_{q_{\sigma}}[\sigma] + \bar{n}_k, \quad \hat{\gamma}_{k,2} = \mathbb{E}_{q_{\alpha}}[\alpha] + k \mathbb{E}_{q_{\sigma}}[\sigma] + \sum_{\ell=k+1}^K \bar{n}_{\ell}, \quad (34)$$

where

$$\text{for } k = 1, \dots, K, \quad \bar{n}_k = \sum_{j=1}^n q_{z_j}(k) \quad (35)$$

corresponds to the weight of cluster k .

VB-E-(α, σ) step

The (α, σ) variational posterior is more complex but has a simple gamma form in the DP ($\sigma = 0$) case. More specifically, we need to compute

$$\hat{s}_1 = s_1 + K - 1, \text{ and } \hat{s}_2 = s_2 - \sum_{k=1}^{K-1} \psi(\hat{\gamma}_{k,2}) - \psi(\hat{\gamma}_{k,1} + \hat{\gamma}_{k,2}) \quad (36)$$

where $\psi(\cdot)$ is the digamma function defined by $\psi(z) = \frac{d}{dz} \log \Gamma(z) = \frac{\Gamma'(z)}{\Gamma(z)}$. When $\sigma = 0$ then q_α is a gamma distribution $\mathcal{G}(\hat{s}_1, \hat{s}_2)$ and $\mathbb{E}_{q_\alpha}[\alpha] = \frac{\hat{s}_1}{\hat{s}_2}$. Otherwise (PY case), $q_{\alpha, \sigma}$ is only identified up to a normalizing constant but the required $\mathbb{E}_{q_\alpha}[\alpha]$ and $\mathbb{E}_{q_\sigma}[\sigma]$ can be computed by importance sampling (see Appendix A.4 for details).

VB-E-Z step

Due to the mean field approximation and the truncation, this step consists in computing (see details in Appendix A.5), for all $j = 1, \dots, n$ and for $k \leq K$,

$$q_{z_j}(k) = \frac{\tilde{q}_j(k)}{\sum_{\ell=1}^K \tilde{q}_j(\ell)}, \quad (37)$$

where $\log \tilde{q}_j(k)$ is defined by

$$\begin{aligned} & -\frac{1}{2} \left\{ \log \left| \frac{\hat{\Psi}_k}{2} \right| - \sum_{i=1}^d \psi \left(\frac{\hat{\nu}_k + (1-i)}{2} \right) + \hat{\nu}_k (y_j - \hat{m}_k)^T \hat{\Psi}_k^{-1} (y_j - \hat{m}_k) + \frac{d}{\hat{\lambda}_k} \right\} + \\ & \psi(\hat{\gamma}_{k,1}) - \psi(\hat{\gamma}_{k,1} + \hat{\gamma}_{k,2}) + \sum_{l=1}^{k-1} \psi(\hat{\gamma}_{l,2}) - \psi(\hat{\gamma}_{l,1} + \hat{\gamma}_{l,2}) + \beta \sum_{i \in \mathcal{N}_j} q_{z_i}(k), \end{aligned} \quad (38)$$

where in the last sum, $\mathcal{N}_j$ represents the neighbours of j . In the above formula, symbols $(\hat{m}_k, \hat{\lambda}_k, \hat{\Psi}_k, \hat{\nu}_k)$ are the variational hyperparameters for $q_{\theta_k^*}$ more specifically defined in the following step and d is the dimension of the data. The advantage of Eq. (37) is that it provides assignment probabilities $q_{z_i}(k)$ and does not require intermediate commitments to hard assignments of the z_j 's. The hard assignments can be postponed to the end if desired to get a segmentation through the following maximum a posteriori (MAP) estimation:

$$\hat{z}_j = \arg \max_{k \in \{1, \dots, K\}} q_{z_j}(k). \quad (39)$$

VB-E- θ^* step

This step is divided into K parts where the computation is similar to that in standard Bayesian finite mixtures with a choice of conjugate prior, here for Gaussian distributions. Hence, for each $k \leq K$, the variational posterior is a Normal-inverse-Wishart distribution defined as

$$q_{\theta_k^*}(\mu_k, \Sigma_k) = \mathcal{NIW}(\mu_k, \Sigma_k; \hat{m}_k, \hat{\lambda}_k, \hat{\Psi}_k, \hat{\nu}_k), \quad (40)$$

where the hyperparameters are updated as follows (see for instance [25])

$$\begin{aligned} \hat{\lambda}_k &= \lambda_k + \bar{n}_k, \quad \hat{\nu}_k = \nu_k + \bar{n}_k, \\ \hat{\Psi}_k &= \Psi_k + S_k + \frac{\lambda_k \bar{n}_k}{\lambda_k + \bar{n}_k} (m_k - \bar{\mu}_k)(m_k - \bar{\mu}_k)^T, \\ \hat{m}_k &= \frac{\lambda_k m_k + \bar{n}_k \bar{\mu}_k}{\lambda_k + \bar{n}_k} = \frac{\lambda_k m_k + \bar{n}_k \bar{\mu}_k}{\hat{\lambda}_k}, \end{aligned} \quad (41)$$

with $\bar{n}_k$ defined in (35) and

$$\begin{aligned} \bar{\mu}_k &= \frac{1}{\bar{n}_k} \sum_{j=1}^n q_{z_j}(k) y_j, \\ S_k &= \sum_{j=1}^n q_{z_j}(k) (y_j - \bar{\mu}_k)(y_j - \bar{\mu}_k)^T. \end{aligned} \quad (42)$$

VB-M steps

The maximization step consists of updating the hyperparameters $\phi = (\beta, s_1, s_2, a, \rho)$, where $\rho = (\rho_1, \dots, \rho_K)$, by maximizing the free energy, if they are not set heuristically:

$$\phi^{(r)} = \arg \max_{\phi} \mathbb{E}_{q_Z^{(r)} q_{\tau}^{(r)} q_{\alpha, \sigma}^{(r)} q_{\theta^*}^{(r)}} [\log p(\mathbf{y}, \mathbf{Z}, \tau, \alpha, \sigma, \theta^*; \phi)] . \quad (43)$$

The VB-M-step can therefore be divided into 3 independent sub-steps as listed below. From the conditional independence of (s_1, s_2, a, ρ) and $(\mathbf{Y}, \mathbf{Z})$ given $(\tau, \alpha, \sigma, \theta^*)$, the VB-M-step writes as in (30) so that the solutions for the VB-M- (s_1, s_2) (in the DP case) and VB-M- ρ steps are straightforward. Only the β step and the M- (s_1, s_2, a) step (in the PY case) are more involved.

VB-M- β : The maximization of (43) with respect to β leads to

$$\beta^{(r)} = \arg \max_{\beta} \mathbb{E}_{q_Z^{(r)} q_{\tau}^{(r)}} [\log p(\mathbf{Z} | \tau; \beta)] . \quad (44)$$

This step does not admit a closed-form solution but can be solved numerically. More details are given in Appendix A.6.

VB-M- (s_1, s_2, a) : This step is straightforward in the DP case ($\sigma = 0$). It can be expressed easily using the fact that both the prior and the variational posterior are Gamma distributions, and using the cross-entropy properties,

$$(s_1, s_2)^{(r)} = \arg \max_{(s_1, s_2)} \mathbb{E}_{q_{\alpha}^{(r)}} [\log p(\alpha; s_1, s_2)] = (\hat{s}_1^{(r)}, \hat{s}_2^{(r)}) \quad (45)$$

where $(\hat{s}_1^{(r)}, \hat{s}_2^{(r)})$ is given in (36). In the more general PY case, we can solve this step numerically using also importance sampling. See Appendix A.7.

VB-M- ρ : This step divides into K sub-steps that involve again cross-entropies,

$$\rho_k^{(r)} = \arg \max_{\rho} \mathbb{E}_{q_{\theta_k^*}^{(r)}} [\log p(\theta_k^*; \rho_k)] = \hat{\rho}_k^{(r)} \quad (46)$$

where $\hat{\rho}_k^{(r)} = (\hat{\lambda}_k^{(r)}, \hat{\nu}_k^{(r)}, \hat{\psi}_k^{(r)}, \hat{m}_k^{(r)})$ is given in Eq. (41).

6 Application to image segmentation

To validate the proposed approach, we consider its application to unsupervised image segmentation as a spatial clustering task. Image segmentation consists of partitioning a digital image into distinct regions that contain pixels with similar properties. Extensive research work has been done in this field using various clustering techniques. In practice, to be meaningful for image analysis and interpretation, the segmented regions should closely relate to depicted objects or features of interest. A number of tasks in image analysis often depends on the reliability of preliminary segments, but an accurate partitioning of an image is still quite challenging. To assess our variational approximation solution, we first provide some results on synthetic grey-level images before testing on real color images for which it is often necessary to consider higher level characteristics (features).

6.1 Experiments on simulated images

In this section, we focus on testing the ability of our variational approximation to handle the typical complications due to the addition of a Potts component, namely the estimation of the Potts interaction parameter β . We simulated grey-level images by adding some Gaussian noise to the realizations of a Potts model with different values of β ($\beta_{\text{true}} \in \{0.6, 0.8, 1\}$) and K ($K_{\text{true}} \in \{5, 7\}$). Given each pair of β_{true} and K_{true} , 100 images of size 64×64 were generated following a Potts model with a second order neighborhood (each pixel has 8 neighbors). The obtained K_{true} different labels were turned into different grey-levels, *eg.* 80, 100, 120, 140, 160, etc. to which we added a centered Gaussian noise with standard deviation equal to $\sqrt{20}$ to get final grey-level images. We then applied our algorithm on each image with a truncation at $K = 30$. The estimated values of β, α, σ were averaged over all 100 simulated images from a given $(\beta_{\text{true}}, K_{\text{true}})$ pair. They are reported in Table 1 with their standard deviations and the frequencies of the numbers of clusters found. The inferred β and K values are generally very close to the true values. In addition, the estimated posterior means of α and σ seem to be quite stable.

$(\beta_{\text{true}}, K_{\text{true}})$	$\bar{\alpha}$	$\text{std}(\alpha)$	$\bar{\sigma}$	$\text{std}(\sigma)$	$\bar{\beta}$	$\text{std}(\beta)$	cluster numbers	frequency (%)
(0.6, 5)	0.96	0.23	0.46	0.19	0.58	0.04	[3, 4, 5 , 6]	[1, 7, 87, 5]
(0.8, 5)	0.90	0.22	0.50	0.16	0.81	0.03	[4, 5 , 6]	[7, 85, 8]
(1.0, 5)	0.98	0.30	0.45	0.18	1.08	0.06	[4, 5 , 6]	[8, 81, 11]
(0.6, 7)	1.09	0.32	0.45	0.28	0.66	0.04	[6, 7 , 8]	[2, 91, 7]
(0.8, 7)	1.00	0.25	0.43	0.21	0.79	0.04	[4, 5, 6, 7 , 8]	[1, 3, 25, 60, 11]
(1.0, 7)	1.03	0.27	0.44	0.21	1.05	0.05	[5, 6, 7 , 8]	[1, 33, 61, 5]

Table 1 Simulated 64×64 images from a Potts model with additional Gaussian noise with varying β and K values (first column). Each model is simulated 100 times. The variational algorithm results are summarized through the parameters means ($\bar{\alpha}$, $\bar{\sigma}$, $\bar{\beta}$) and their standard deviations. The numbers of clusters found are also given with their frequencies (most frequent number in bold characters).

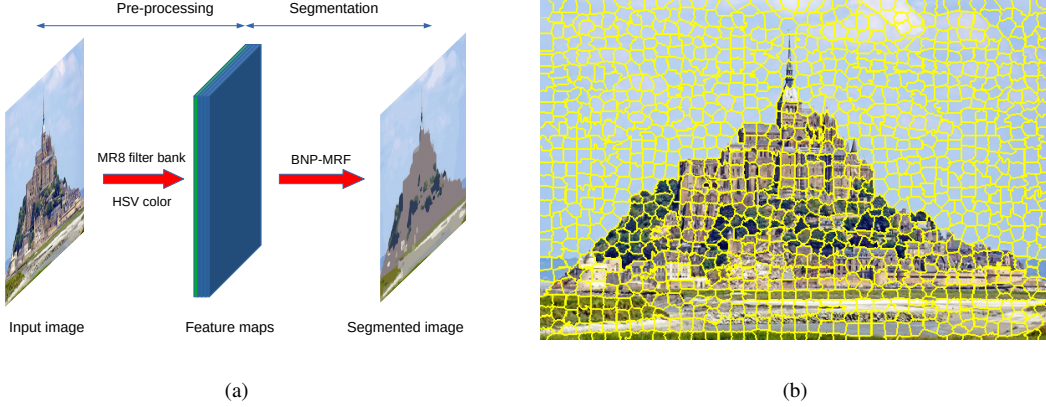


Fig. 3 Pre-processing steps: (a) color and texture feature extraction, each pixel is associated to 3 color indices (green slice in the middle image) and 8 texture indices (blue slice) ; (b) example of a super-pixel pre-segmentation.

6.2 Feature extraction for image segmentation

The goal of segmentation is to group pixels that are similar usually in terms of color and texture. The evaluation of this similarity can be based on the comparison of quantitative characteristics or features. The color and texture features in a natural image are often very complex. For our experiments, we mainly focus on two special types of features based on the HSV (Hue, Saturation, Value) color space and the maximum response (MR) filter bank. The HSV color space is often used in natural image analysis because it corresponds better to how people experience color than the RGB color space does. Regarding the texture information, we shall consider the MR8 filter bank [39], which consists of 38 filters but only 8 filter responses. More precisely, the MR8 filter bank contains filters at multiple orientations but their outputs are compressed by recording only the maximum filter response across all orientations. This achieves rotation invariance. It follows that each pixel can be associated to three color indices and eight texture indices resulting in 11 features as summarized in Figure 3(a). Furthermore, the images are presegmented into super-pixels that group pixels similar in color and other low-level properties [1]. More specifically, pixels are grouped to form a single super-pixel. The super-pixel representation is a frequently used techniques to pre-group pixels in image processing. In this respect, super-pixels are regarded as more natural entities that allow reducing the number of observations drastically for running clustering algorithms. In all our experiments, each image is pre-segmented into approximately 1 000 super-pixels using the SLIC algorithm proposed in [1]. For illustration, Figure 3(b) shows a super-pixel segmentation. Finally, we compute the 11-dimensional feature vectors at the super-pixel level by considering the centroid of each super-pixel and taking the average of features over the pixels in the super-pixel. The entire segmentation procedure is summarized in Algorithm 1.

6.3 Berkeley Segmentation Data Set

Previous empirical studies on similar annotated image data [36,32] have shown that segment sizes in a set of manually segmented images follow a power law distribution well modeled by a Pitman–Yor process. The Berkeley Segmentation Data Set 500 (BSDS500) contains 500 images associated each with several manual segmentations, for a total number of 2696 manually segmented images. From these segmentations, we computed for

Algorithm 1: Summary of the image segmentation procedure.**Input:** An input color image.**Output:** A segmented image.**Procedure:**

- Use the SLIC algorithm to form an over-segmentation [1].
- Compute HSV color and texture features at super-pixel level [39].
- Partition the aggregated features by BNP-MRF.

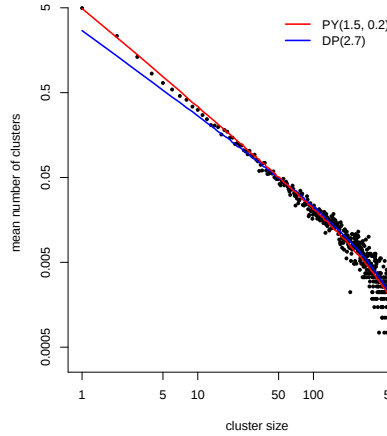
Return: Multiple segmented regions.

Fig. 4 Empirical segment sizes (black dots) computed over the 2696 manually segmented images of BSDS500. The curves represent theoretical segment size distributions obtained with 10^5 simulations respectively from Pitman–Yor ($\alpha = 1.5$, $\sigma = 0.2$, red curve) and Dirichlet processes ($\alpha = 2.7$, blue curve).

each of them the segment sizes evaluated in number of super-pixels to obtain the empirical segment size distribution shown in Figure 4. To check the adequation with the Pitman–Yor and Dirichlet process size distributions, we simulated 10^5 samples of size $n = 1000$ of both processes. These samples were then used to evaluate the theoretical size distributions for given values of the parameters α and σ . A sample size of 1000 was chosen so as to match the image sizes in terms of super-pixels. The resulting curves are shown, after additional smoothing, in Figure 4. They exhibit similar shapes as in [40] (Figure 2) which interestingly also shows the deviation from the usual asymptotic formulas (eq. (6) and (8) in [40]) when n tends to infinity. In Figure 4, the plotted curves correspond to the best fits we could find for the Pitman–Yor and Dirichlet processes. They show that the Pitman–Yor process fits better the empirical distribution for small segment sizes with a steeper slope than that of the DP (see the discussion in Section 4.3). For large segment sizes the empirical distribution is not precise enough to conclude on a real difference between the two processes. In terms of segmentation quality, both options would probably be acceptable, but we report below mainly results with the Pitman–Yor process for its extra flexibility.

To quantify the performance of our segmentation algorithm, numerical experiments were conducted on a subset of images selected from the BSDS500 data set already studied by [3, 12], which provides multiple human annotated segments as many ground truths for each image. The considered subset consists of 154 images as listed in Tables 1 and 2 in [12].

In the literature, a standard measure for comparing a test segmentation to another is the rand index (RI) [30]. The RI is one when two segmentations are exactly the same. However, when having for one image a set of ground truths which do not completely agree, the probabilistic rand index (PRI) [38] is preferable. Given a set of ground truths $\mathcal{S} = (S_1, \dots, S_T)$, the PRI is defined as follows:

$$\text{PRI}(S_{\text{test}}, \mathcal{S}) = \frac{2}{n(n-1)} \sum_{i < j} [c_{ij}p_{ij} + (1 - c_{ij})(1 - p_{ij})] \quad (47)$$

where $c_{ij} = 1$ if pixels i and j belong to the same segment in S_{test} and $c_{ij} = 0$ otherwise, n is the number of image pixels and p_{ij} is the probability of two pixels i and j having the same label, *i.e.*, the fraction of all available ground truths in $\mathcal{S}$ where pixels i and j belong to the same segment. In fact, it can be shown that Eq. (47) is simply the mean of the RI computed between each pair (S_{test}, S_k) , namely $\frac{1}{T} \sum_{k=1}^T \text{RI}(S_{\text{test}}, S_k)$. By definition, the PRI always takes values in $[0, 1]$, where 0 means that S_{test} and $(S_1, \dots, S_T)$ have no similarities and 1 means

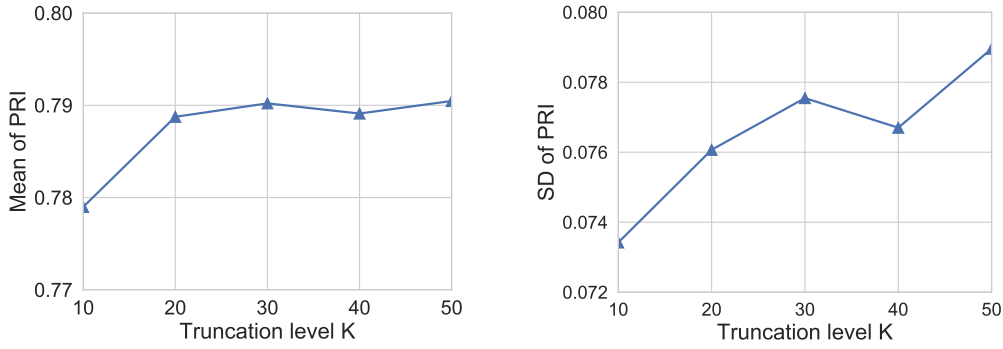


Fig. 5 PY-MRF mixture model: Mean and standard deviation of the PRI score over the considered subset of the BSDS500 data set as a function of the truncation level K .

Proposed model		Results given in [12]			
PRI (%)	PY-MRF	DPM	iHMRF	MRF-PYP	Graph Cuts
Mean	79.05	74.15	75.50	76.49	76.10
Median	80.62	75.49	76.89	78.08	77.59
St. Dev.	7.9	8.4	8.2	7.9	8.3

Table 2 Performance comparison: summary statistics of the PRI score over the 154 images from BSDS500 studied by [13] for our PY-MRF mixture model and the approaches tested in [13, 12].

all segments are identical. The larger the PRI, the better. In practice, PRI values are often reported as percentages in $[0, 100]$.

Our approach has been tested on the considered subset of the BSDS500 and the summary statistics of the PRI score are shown in Figure 5 as a function of the truncation level K for the PY-MRF case. Similar results were observed for the DP. It appears that for $K \geq 30$, the global performance does not change much and is satisfying with respect to existing results in the literature. We compared our best results with those reported in [12]. Table 2 shows that our approach outperforms the existing results. The improvement in PRI may appear overall small but it can be assessed by visualizing original images and their segmentations. We show in Figure 6 segmentation results for four images. The main differences between the non spatial PY and PY-MRF mixture models can be visualized for the first image in the ground and water which are segmented in the latter case into a smaller number of regions whose shapes are in addition smoother. This is typical of more spatial interaction in the clustering process. Similarly, the same phenomenon is also visible in the peak part of the second image, in the sky and grass parts of the third image and in the plant parts of the fourth image.

We also examined, for the PY-MRF mixture model with $K = 50$, the values of the expected α , σ and β for each of the 154 segmented images presented in Figure 7 as scatter plots (one point per image). Recall from Section 5.2 that α and σ are elements of the parameters Θ while β is considered as a hyperparameter from ϕ . Figure 7 shows also the correlations (across the 154 images) between the expected values of α , σ and β . It appears that the estimated σ values are most of the time smaller than 0.5 and sometimes closer to 0 with some anti-correlation with respect to α values. In contrast, β values appear quite independent from α or σ .

In terms of pure PRI performance, the BSDS500 data set is not an easy example because the ground truth segmentations are labeled manually by humans and are sometimes quite subjective and inconsistent across users. However, this example allows comparison of methods and visualization. Two interesting findings are that the choice of K does not seem to be too sensitive as soon as K is large enough, and there seems to be some correlation between α and σ while β is rather independent of the latest. Further analysis would be needed to confirm these properties but in practice, they could be used to guide the segmentations into more or less spatially smooth versions without risking to eliminate too small segments.

The code used for these experiments is available at <https://team.inria.fr/mistis/software/> under the name BNP-MRF.

7 Conclusion and perspectives

In this paper, we proposed a general scheme to build BNP priors that can model dependencies through the addition of a Markov random field term. In contrast to other existing attempts that reduce to spatially constrained



Fig. 6 Segmentation results for four images from the BSDS500 data set. From left to right, columns show respectively, the original images, the segmentation results with the PY and PY-MRF mixture models.

standard BNP priors such as [12, 13], our proposal leads to proper spatial priors. Our construction is based on the stick-breaking representation and was illustrated starting from the Dirichlet and Pitman–Yor processes, although this approach could be extended to other forms of BNP priors admitting a stick-breaking representation such as Gibbs-type priors. The stick-breaking representation was further exploited to derive clustering properties of the model and to provide a variational inference algorithm. In addition to the usual BNP parameters, an estimation of the Markov interaction parameter β was proposed. The variational approximation chosen was based on a standard truncation but it would be interesting to investigate other approximations, *e.g.* [41]. Also the variational algorithm is greatly simplified for standard stick-breaking representations (*e.g.* DP and PY) with independent weight variables. Nevertheless, it would be interesting to investigate more general stick-breaking representations possibly using some MCMC counterpart for estimation.

The approach was illustrated on a challenging unsupervised image segmentation task with good results with respect to the literature, but the proposed scheme is quite flexible and can be used in more general settings including community detection or disease mapping in epidemiology.

A Appendix

Proofs of propositions 1 and 2 are given in the two first sections. Details on the VBEM steps, to complete the previous developments when necessary, follow.

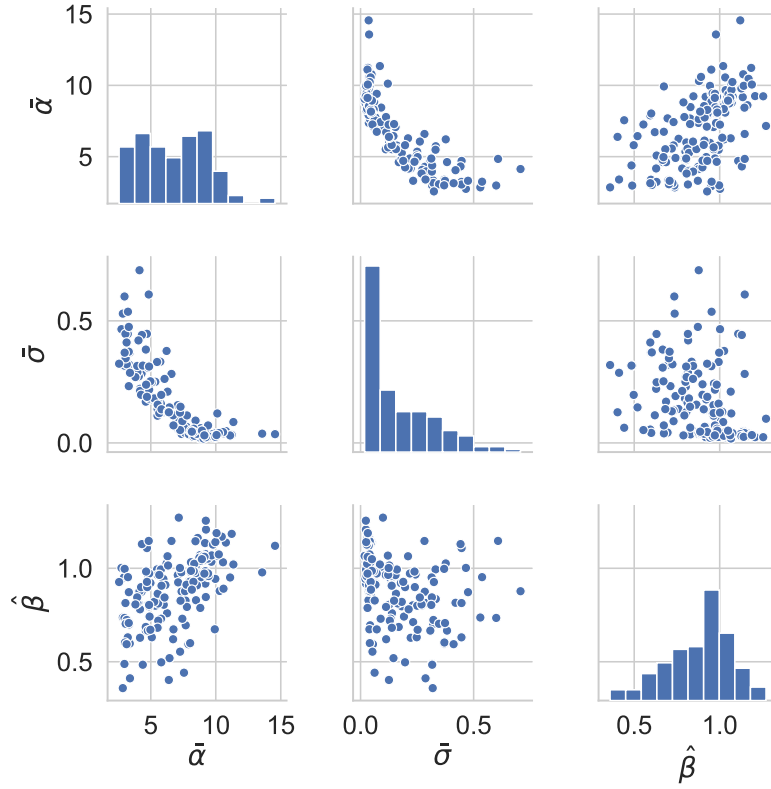


Fig. 7 Estimated parameter values $(\hat{\alpha}, \hat{\sigma})$ obtained from **VB-E** steps and $\hat{\beta}$ obtained from a **VB-M** step using the PY-MRF model with truncation level $K = 50$, on the 154 images from the Berkeley benchmark.

A.1 Proof of Proposition 1

Integrating out π from distribution (12) and noticing that the Markov term does not depend on π , leads to

$$p(z_{1:n+1}; \beta) \propto p(z_{1:n+1}; \beta = 0) e^{\beta H(z_{1:n+1})}$$

where $H(z_{1:n+1}) = \sum_{i \sim j} \delta_{z_i = z_j}$ is counting the number of homogeneous edges. It follows that

$$p(z_{n+1} | z_{1:n}; \beta) \propto p(z_{n+1} | z_{1:n}; \beta = 0) e^{\beta H(z_{1:n+1})}.$$

When $\beta = 0$, the model (12) reduces to standard Gibbs-type priors for which the quantities above have been given in (23). Using (23) and enumerating how z_{n+1} can affect the number of homogeneous edges, it comes that: for $z_{n+1} \notin \{z_1, \dots, z_n\}$,

$$p(z_{n+1} | z_{1:n}; \beta) \propto \frac{V_{n+1, K_{n+1}}}{V_{n, K_n}} e^{\beta H(z_{1:n})}$$

and for $\ell \in \{z_1 \dots z_n\}$,

$$p(z_{n+1} = \ell | z_{1:n}; \beta) \propto \frac{V_{n+1, K_n}}{V_{n, K_n}} (n_\ell - \sigma) e^{\beta \tilde{n}_\ell \delta_{N_{n+1}}(\ell)} e^{\beta H(z_{1:n})}.$$

The result follows by normalizing the quantities above, using (21) and noticing that

$$\sum_{\ell \in \{z_1 \dots z_n\}} (n_\ell - \sigma) e^{\beta \tilde{n}_\ell \delta_{N_{n+1}}(\ell)} = \eta_{n+1} + n - \sigma K_n.$$

■

A.2 Proof of Proposition 2

Consider the DP first. Let $D_n = \delta_{\theta_n \text{ is new} | \theta_{1:n-1}}$ be the Bernoulli random variable equal to one when θ_n is a fresh draw, which for the DP-MRF happens with probability

$$\frac{\alpha}{\alpha + n - 1 + \boldsymbol{\eta}_n(0, \beta)}.$$

Then $K_n = \sum_{i=1}^n D_i$, so we have for the DP-MRF

$$\mathbb{E}[K_n] = \sum_{i=0}^{n-1} \mathbb{E} \left[\frac{\alpha}{\alpha + i + \boldsymbol{\eta}_{i+1}(0, \beta)} \right]. \quad (48)$$

By definition of the maximal degree D of the graph, the multiplicity $\tilde{n}_\ell$ is at most D for any ℓ , hence we have the following upper bound

$$\sum_{\ell \in \mathcal{Z}_{\mathcal{N}_{i+1}}} n_\ell (e^{\beta \tilde{n}_\ell} - 1) \leq \sum_{\ell \in \mathcal{Z}_{\mathcal{N}_{i+1}}} n_\ell (e^{D\beta} - 1) \leq i(e^{D\beta} - 1). \quad (49)$$

Plugging (49) into (48) provides the desired inequality

$$\begin{aligned} \mathbb{E}[K_n] &\geq \sum_{i=0}^{n-1} \frac{\alpha}{\alpha + i e^{D\beta}} \\ &= \frac{\alpha}{e^{D\beta}} \sum_{i=0}^{n-1} \frac{1}{i + \alpha e^{-D\beta}} \gtrsim \frac{\alpha}{e^{D\beta}} \log n. \end{aligned}$$

Turning to the Pitman–Yor process, we follow the proof technique of [37]. The prior expectation of the number of clusters satisfies the following recursion

$$\mathbb{E}[K_{n+1}] = \mathbb{E}[K_n] + \mathbb{E} \left[\frac{\alpha + \sigma K_n}{\alpha + n + \boldsymbol{\eta}_{n+1}} \right].$$

Assuming here that $\alpha \geq 0$ and using the lower bound (49) yields

$$\mathbb{E}[K_{n+1}] \geq \mathbb{E}[K_n] \left(1 + \frac{\sigma}{\alpha + n e^{D\beta}} \right).$$

By induction, and using $K_1 = 1$,

$$\begin{aligned} \mathbb{E}[K_n] &\geq \prod_{i=1}^{n-1} \left(1 + \frac{\sigma}{\alpha + i e^{D\beta}} \right) = \exp \left(\sum_{i=1}^{n-1} \log \left(1 + \frac{\sigma}{\alpha + i e^{D\beta}} \right) \right) \\ &\simeq \exp \left(\sum_{i=1}^{n-1} \frac{\sigma}{\alpha + i e^{D\beta}} \right) \simeq \exp \left(\frac{\sigma}{e^{D\beta}} \log n \right) = n^{\sigma e^{-D\beta}}, \end{aligned}$$

where $a_n \simeq b_n$ means a_n/b_n converges to some positive constant when $n \rightarrow \infty$. The desired inequality for PY is obtained by writing this constant equal to c . $\blacksquare$

A.3 VB-E- τ step

$$\begin{aligned} q_{\tau_k}(\tau_k) &\propto \exp \left(\mathbb{E}_{q_{\alpha, \sigma}} [\log p(\tau_k | \alpha, \sigma)] + \sum_{j=1}^n \mathbb{E}_{q_{z_j} q_{\tau \setminus \{k\}}} [\log \pi_{z_j}(\tau)] \right) \\ &\propto \exp \left(-\mathbb{E}_{q_{\alpha, \sigma}} [\sigma] \log \tau_k + (\mathbb{E}_{q_{\alpha, \sigma}} [\alpha] + k \mathbb{E}_{q_{\alpha, \sigma}} [\sigma] - 1) \log(1 - \tau_k) \right. \\ &\quad \left. + \sum_{j=1}^n \sum_{l=k+1}^K q_{z_j}(l) \log(1 - \tau_k) + \sum_{j=1}^n q_{z_j}(k) \log(\tau_k) \right). \end{aligned} \quad (50)$$

Considering the terms involving τ_k , we recognize the beta distribution $\mathcal{B}(\tau_k; \hat{\gamma}_{k,1}, \hat{\gamma}_{k,2})$ specified in (33). It follows the expressions of the following quantities,

$$\begin{aligned} \mathbb{E}_{q_{\tau_k}} [\log(\tau_k)] &= \psi(\hat{\gamma}_{k,1}) - \psi(\hat{\gamma}_{k,1} + \hat{\gamma}_{k,2}) \\ \mathbb{E}_{q_{\tau_k}} [\log(1 - \tau_k)] &= \psi(\hat{\gamma}_{k,2}) - \psi(\hat{\gamma}_{k,1} + \hat{\gamma}_{k,2}) \end{aligned} \quad (51)$$

where $\psi(\cdot)$ is the digamma function defined in Section 5.2.

A.4 VB-E-(α, σ) step

In the PY case, $q_{\alpha, \sigma}(\alpha, \sigma)$ is proportional to

$$\begin{aligned} \tilde{q}_{\alpha, \sigma}(\alpha, \sigma) &\propto p(\alpha, \sigma; s_1, s_2, a) \exp \left(\sum_{k=1}^{K-1} \mathbb{E}_{q_{\tau_k}} [\log p(\tau_k | \alpha, \sigma)] \right) \\ &\propto p(\alpha, \sigma; s_1, s_2, a) \prod_{k=1}^{K-1} B(1 - \sigma, \alpha + k\sigma)^{-1} \times \\ &\exp \left(-\sigma \left(\sum_{k=1}^{K-1} \mathbb{E}_{q_{\tau_k}} [\log \tau_k] - \sum_{k=1}^{K-1} k \mathbb{E}_{q_{\tau_k}} [\log(1 - \tau_k)] \right) + (\alpha - 1) \sum_{k=1}^{K-1} \mathbb{E}_{q_{\tau_k}} [\log(1 - \tau_k)] \right), \end{aligned}$$

where $B(a, b) := \frac{\Gamma(a)\Gamma(b)}{\Gamma(a+b)}$ represents the beta function. The last term simplifies into

$$\begin{aligned} \prod_{k=1}^{K-1} B(1 - \sigma, \alpha + k\sigma)^{-1} &= \frac{\Gamma(\alpha)}{\Gamma(1 - \sigma)^{K-1} \Gamma(\alpha + (K-1)\sigma)} \\ &\times \prod_{k=1}^{K-1} (\alpha + (k-1)\sigma). \end{aligned} \quad (52)$$

The difficulty is that except in the DP-MRF case, the normalizing constant for $\tilde{q}_{\alpha, \sigma}(\alpha, \sigma)$ is not tractable. Nevertheless, to carry out the VBEM algorithm, we do not need the full $q_{\alpha, \sigma}$ distribution but only the means $\mathbb{E}_{q_{\sigma}}[\sigma]$ and $\mathbb{E}_{q_{\alpha}}[\alpha]$. One solution is therefore to use importance sampling or MCMC to compute these expectations via Monte Carlo sums. Using the prior on (α, σ) given in (31), it comes that

$$\begin{aligned} \tilde{q}_{\alpha, \sigma}(\alpha, \sigma) &= \mathcal{G}(\alpha + \sigma; \hat{s}_1, \hat{s}_2) e^{-\sigma \xi} \\ &p(\sigma; a) \frac{\Gamma(\alpha)}{\Gamma(1 - \sigma)^{K-1} \Gamma(\alpha + (K-1)\sigma)} \prod_{k=1}^{K-1} \left(\frac{\alpha + (k-1)\sigma}{\alpha + \sigma} \right) \end{aligned} \quad (53)$$

where $\mathcal{G}(\alpha + \sigma; \hat{s}_1, \hat{s}_2)$ is the pdf of a σ -shifted gamma distribution with $\hat{s}_1, \hat{s}_2$ given in (36). The parameter ξ is defined as

$$\xi = \sum_{k=1}^{K-1} \mathbb{E}_{q_{\tau_k}} [\log \tau_k] - \sum_{k=1}^{K-1} (k-1) \mathbb{E}_{q_{\tau_k}} [\log(1 - \tau_k)], \quad (54)$$

which can be computed using (51). We propose then to use as importance distribution $\nu(\alpha, \sigma) = \mathcal{G}(\alpha + \sigma; \hat{s}_1, \hat{s}_2) p(\sigma; a)$ with $p(\sigma; a)$ the uniform distribution on $[0, 1]$, $\mathcal{U}_{[0, 1]}(\sigma)$. It comes an expression for the importance weights,

$$\begin{aligned} W(\alpha, \sigma) &= \frac{\tilde{q}_{\alpha, \sigma}(\alpha, \sigma)}{\nu(\alpha, \sigma)} \\ &= e^{-\sigma \xi} \frac{\Gamma(\alpha)}{\Gamma(1 - \sigma)^{K-1} \Gamma(\alpha + (K-1)\sigma)} \prod_{k=1}^{K-1} \left(\frac{\alpha + (k-1)\sigma}{\alpha + \sigma} \right). \end{aligned} \quad (55)$$

The importance sampling scheme then consists of:

- For $i = 1$ to M , simulate first independently σ_i from $\mathcal{U}_{[0, 1]}(\sigma)$ and then simulate conditionnaly α_i using the σ_i -shifted gamma $SG(\sigma_i, \hat{s}_1, \hat{s}_2)$. This later simulation is easily obtained by simulating a standard $\mathcal{G}(\hat{s}_1, \hat{s}_2)$ and then subtracting σ_i to the result.
- Compute the importance weights $w_i = W(\alpha_i, \sigma_i)$.
- Approximate the means $\mathbb{E}_{q_{\sigma}}[\sigma] \approx \frac{\sum_{i=1}^M w_i \sigma_i}{\sum_{i=1}^M w_i}$ and $\mathbb{E}_{q_{\alpha}}[\alpha] \approx \frac{\sum_{i=1}^M w_i \alpha_i}{\sum_{i=1}^M w_i}$

Note that this complication is due to the PY. In the DP-MRF case, the E- α step is considerably simpler as it reduces to computing the approximate posterior expectation of α , namely $\mathbb{E}_{q_{\alpha}}[\alpha] = \hat{s}_1 / \hat{s}_2$.

A.5 VB-E-Z step

The VB-E-Z is divided into n steps. Since we assume $q_{z_j}(z_j) = 0$ for $z_j > K$, it is only necessary to compute the distributions for $z_j \leq K$, namely

$$q_{z_j}(z_j) \propto \exp \left(\mathbb{E}_{q_{\theta_{z_j}^*}} [\log p(y_j | \theta_{z_j}^*)] + \mathbb{E}_{q_{\tau}} [\log \pi_{z_j}(\tau)] + \beta \sum_{i \in \mathcal{N}_j} q_{z_i}(z_j) \right), \quad (56)$$

where for $z_j = k$,

$$\mathbb{E}_{q_{\tau}} [\log \pi_k(\tau)] = \mathbb{E}_{q_{\tau_k}} [\log \tau_k] + \sum_{l=1}^{k-1} \mathbb{E}_{q_{\tau_l}} [\log(1 - \tau_l)]. \quad (57)$$

The term $\mathbb{E}_{q_{\theta^*}^{z_j}}[\log p(y_j | \theta_{z_j}^*)]$ is computed using the fact that $q_{\theta^*}^{z_j}$ is a Normal-inverse-Wishart distribution as described in Eq. (40) of the next VB-E- θ^* step, namely

$$\begin{aligned} q_{\theta^*}^{z_j}(\mu_k, \Sigma_k) &= \mathcal{N}\mathcal{IW}(\mu_k, \Sigma_k; \hat{m}_k, \hat{\lambda}_k, \hat{\Psi}_k, \hat{\nu}_k) \\ &= \mathcal{N}\left(\mu_k; \hat{m}_k, \frac{\Sigma_k}{\hat{\lambda}_k}\right) \mathcal{IW}(\Sigma_k; \hat{\Psi}_k, \hat{\nu}_k). \end{aligned} \quad (58)$$

It comes out that (d is the dimension of y_j)

$$\begin{aligned} \mathbb{E}_{q_{\theta^*}^{z_j}}[\log p(y_j | \theta_k^*)] &= -\frac{d}{2} \log 2\pi - \frac{1}{2} \mathbb{E}_{q_{\theta^*}^{z_j}}[\log |\Sigma_k|] \\ &\quad - \frac{1}{2} \mathbb{E}_{q_{\theta^*}^{z_j}}[(y_j - \mu_k)^T \Sigma_k^{-1} (y_j - \mu_k)]. \end{aligned} \quad (59)$$

Using the decomposition (cf. Eq. (58)) and the fact that Σ_k^{-1} follows a Wishart distribution, it comes out that

$$\begin{aligned} \mathbb{E}_{q_{\theta^*}^{z_j}}[\log |\Sigma_k|] &= \mathbb{E}_{q_{\Sigma_k}}[\log |\Sigma_k|] = -\mathbb{E}_{q_{\Sigma_k}}[\log |\Sigma_k^{-1}|] \\ &= -\sum_{i=1}^d \psi\left(\frac{\hat{\nu}_k + (1-i)}{2}\right) + \log \left| \frac{\hat{\Psi}_k}{2} \right|, \end{aligned} \quad (60)$$

Then, one has

$$\begin{aligned} \mathbb{E}_{q_{\theta^*}^{z_j}}[(y_j - \mu_k)^T \Sigma_k^{-1} (y_j - \mu_k)] &= \mathbb{E}_{q_{\Sigma_k}}[(y_j - \hat{m}_k)^T \Sigma_k^{-1} (y_j - \hat{m}_k) + \text{Tr}(\Sigma_k^{-1} \Sigma_k / \hat{\lambda}_k)] \\ &= \hat{\nu}_k (y_j - \hat{m}_k)^T \hat{\Psi}_k^{-1} (y_j - \hat{m}_k) + \frac{d}{\hat{\lambda}_k}. \end{aligned} \quad (61)$$

Plugging in all of the above expressions back into Eq. (56) yields, for $z_j = k$ and $k = 1, \dots, K$, $q_{z_j}(k) \propto \tilde{q}_j(k)$ with $\tilde{q}_j(k)$ given in (38).

A.6 VB-M- β step

This step does not admit a closed-form expression but can be solved numerically. The maximization in β admits a unique solution. Indeed, it is equivalent to solve the following equation:

$$\begin{aligned} \hat{\beta} &= \arg \max_{\beta} \mathbb{E}_{q_{\mathbf{z}} q_{\boldsymbol{\tau}}} [\log p(\mathbf{z} | \boldsymbol{\tau}; \beta)] \\ &= \arg \max_{\beta} \mathbb{E}_{q_{\mathbf{z}} q_{\boldsymbol{\tau}}} [V(\mathbf{z}; \boldsymbol{\tau}, \beta)] - \mathbb{E}_{q_{\boldsymbol{\tau}}} [\log C(\boldsymbol{\tau}, \beta)]. \end{aligned} \quad (62)$$

where C denotes the normalizing constant that depends on $\boldsymbol{\tau}$ and β . Denoting the gradient vector and Hessian matrix respectively by ∇_{β} and ∇_{β}^2 , it comes that

$$\begin{aligned} \nabla_{\beta} \mathbb{E}_{q_{\mathbf{z}} q_{\boldsymbol{\tau}}} [\log p(\mathbf{z} | \boldsymbol{\tau}; \beta)] &= \mathbb{E}_{q_{\mathbf{z}} q_{\boldsymbol{\tau}}} [\nabla_{\beta} V(\mathbf{z} | \boldsymbol{\tau}; \beta)] - \mathbb{E}_{p(\mathbf{z} | \boldsymbol{\tau}; \beta) q_{\boldsymbol{\tau}}} [\nabla_{\beta} V(\mathbf{z} | \boldsymbol{\tau}; \beta)], \\ \nabla_{\beta}^2 \mathbb{E}_{q_{\mathbf{z}} q_{\boldsymbol{\tau}}} [\log p(\mathbf{z} | \boldsymbol{\tau}; \beta)] &= \mathbb{E}_{q_{\mathbf{z}} q_{\boldsymbol{\tau}}} [\nabla_{\beta}^2 V(\mathbf{z} | \boldsymbol{\tau}; \beta)] - \mathbb{E}_{p(\mathbf{z} | \boldsymbol{\tau}; \beta) q_{\boldsymbol{\tau}}} [\nabla_{\beta}^2 V(\mathbf{z} | \boldsymbol{\tau}; \beta)] \\ &\quad - \mathbb{E}_{q_{\boldsymbol{\tau}}} [\text{Var}_{p(\mathbf{z} | \boldsymbol{\tau}; \beta)} [\nabla_{\beta} V(\mathbf{z} | \boldsymbol{\tau}; \beta)]]. \end{aligned} \quad (63)$$

It follows that whenever $V(\mathbf{z} | \boldsymbol{\tau}; \beta)$ is linear in β , $\nabla_{\beta}^2 V(\mathbf{z} | \boldsymbol{\tau}; \beta)$ is zero, the Hessian matrix is negative semidefinite and the function to optimize is concave.

Unfortunately, due to the intractable normalizing constant C , it is necessary to evaluate the terms involving $p(\mathbf{z} | \boldsymbol{\tau}; \beta)$ in an approximate manner. A natural approach is to use a mean-field-like approximation that consists of replacing all neighbors in the interaction term by non-random values. Thus, the Potts model is approximated by

$$p(\mathbf{z} | \boldsymbol{\tau}; \beta) \simeq \prod_{j=1}^n p_{z_j}^{\text{MF}}(z_j | \boldsymbol{\tau}; \beta), \quad (64)$$

with $p_{z_j}^{\text{MF}}(z_j | \boldsymbol{\tau}; \beta)$ defined as

$$p_{z_j}^{\text{MF}}(z_j = k | \boldsymbol{\tau}; \beta) \propto \exp \left(\log \pi_k(\boldsymbol{\tau}) + \beta \sum_{i \in \mathcal{N}_j} q_{z_i}(k) \right). \quad (65)$$

This approximation induced by the posterior variational approximation has been proposed in [9, 18] and also exploited in [10]. It thus follows that

$$\begin{aligned}\mathbb{E}_{q_{\mathbf{z}}q_{\boldsymbol{\tau}}}[\nabla_{\beta}V(\mathbf{z} \mid \boldsymbol{\tau}; \beta)] &= \sum_{k=1}^K \sum_{i \in \mathcal{N}_j} q_{z_j}(k) q_{z_i}(k) \\ \mathbb{E}_{p(\mathbf{z} \mid \boldsymbol{\tau}; \beta)q_{\boldsymbol{\tau}}}[\nabla_{\beta}V(\mathbf{z} \mid \boldsymbol{\tau}; \beta)] &\simeq \mathbb{E}_{q_{\boldsymbol{\tau}}} \left[\sum_{k=1}^K \sum_{i \in \mathcal{N}_j} p_{z_j}^{\text{MF}}(k \mid \boldsymbol{\tau}; \beta) p_{z_i}^{\text{MF}}(k \mid \boldsymbol{\tau}; \beta) \right].\end{aligned}\quad (66)$$

The additional difficulty is that we have to compute the expectation with respect to each q_{τ_k} . This can be done using simulations. To avoid the Monte Carlo sum, we can use instead of (65) another approximation where the random $\boldsymbol{\tau}$ is replaced by a set of fixed values $\tilde{\boldsymbol{\tau}}$ given by

$$\tilde{\tau}_k = \mathbb{E}_{q_{\tau_k}}[\tau_k] = \frac{\hat{\gamma}_{k,1}}{\hat{\gamma}_{k,1} + \hat{\gamma}_{k,2}}. \quad (67)$$

Thus, Eq. (65) turns into

$$p_{z_j}^{\text{MF}}(z_j = k; \beta) \propto \exp \left(\log \pi_k(\tilde{\boldsymbol{\tau}}) + \beta \sum_{i \in \mathcal{N}_j} q_{z_i}(k) \right), \quad (68)$$

where $\log \pi_k(\tilde{\boldsymbol{\tau}}) = \log \tilde{\tau}_k + \sum_{l=1}^{k-1} \log(1 - \tilde{\tau}_l)$. Similarly, it follows that

$$\mathbb{E}_{p(\mathbf{z} \mid \boldsymbol{\tau}; \beta)q_{\boldsymbol{\tau}}}[\nabla_{\beta}V(\mathbf{z} \mid \boldsymbol{\tau}; \beta)] \simeq \sum_{k=1}^K \sum_{i \in \mathcal{N}_j} p_{z_j}^{\text{MF}}(k; \beta) p_{z_i}^{\text{MF}}(k; \beta). \quad (69)$$

By setting them equal to each other and solving this equation for β over an interval, say $[0, 10]$, one obtains an updated value $\hat{\beta}$.

A.7 VB-M- (s_1, s_2, a) step

For the sake of simplicity, we use for σ a uniform prior so that parameter a does not have to be taken into account. With the previous choice of priors on α and σ for PY-MRF mixture models, this VB-M-step becomes

$$(s_1, s_2)^{(r)} = \arg \max_{(s_1, s_2)} \mathbb{E}_{q_{\alpha, \sigma}^{(r)}} [\log p(\alpha \mid \sigma; s_1, s_2)] \quad (70)$$

But the issue is now that the precise form of $q_{\alpha, \sigma}^{(r)}$ is not known. We can use again importance sampling for the optimization.

More specifically, these steps do not admit an explicit closed-form expression but can be solved numerically using gradient ascent schemes. Indeed, for s_1, s_2 , it is equivalent to solve

$$\begin{aligned}\nabla_{s_1} \mathbb{E}_{q_{\alpha}} [\log p(\alpha \mid \sigma; s_1, s_2)] &= \mathbb{E}_{q_{\alpha}} [\nabla_{s_1} \log p(\alpha \mid \sigma; s_1, s_2)] \\ &= \log s_2 + \mathbb{E}_{q_{\alpha, \sigma}} [\log(\alpha + \sigma)] - \Psi(s_1) \\ &= 0 \\ \nabla_{s_2} \mathbb{E}_{q_{\alpha}} [\log p(\alpha \mid \sigma; s_1, s_2)] &= \mathbb{E}_{q_{\alpha}} [\nabla_{s_2} \log p(\alpha \mid \sigma; s_1, s_2)] \\ &= \frac{s_1}{s_2} - \mathbb{E}_{q_{\alpha}} [\alpha] - \mathbb{E}_{q_{\sigma}} [\sigma] \\ &= 0.\end{aligned}\quad (71)$$

As before, when $\sigma = 0$, this step simplifies into $s_1^{(r)} = \hat{s}_1^{(r)}$ and $s_2^{(r)} = \hat{s}_2^{(r)}$, namely to the DP-MRF case.

A.8 Initialization of the VBEM algorithm

An important question which is not often addressed in details is how to initialize the VBEM algorithm. In contrast to the standard EM that we can equivalently start with an initial E-step or an initial M-step, VBEM requires several steps to be initialized depending on the complexity of the model.

In the present work, we propose the following procedure for initializing the VBEM algorithm. First, we set values for s_1 and s_2 , which are taken in our experiments to $s_1 = 1$ and $s_2 = 200/K$. From this, we can initialize the VB-E- (α, σ) step by setting $\mathbb{E}[\alpha] = s_1/s_2$ and $\mathbb{E}[\sigma] = 0$. It is then required to set values for the other hyperparameters. The interaction parameter β can be initialized to $\beta = 0$ assuming no initial spatial interaction and for the ρ_k 's defining the Normal-inverse-Wishart priors, we suggest to use for all k , $m_k = \mathbf{0}$, $\Psi_k = \mathbf{1} \cdot 10^3$, $\nu_k = d$ and $\lambda_k = 1$ in order to start with a large variance for the μ_k 's.

In addition, we need to initialize the cluster assignments which correspond to the VB-E- $\mathbf{Z}$ step. Several approaches are possible depending on the available information and model. A common way often used in image segmentation is to initialize the $q_{\mathbf{Z}}(z_j)$'s randomly or using an initial segmentation coming either from preliminary information or from a simpler non spatial clustering procedure, e.g. *k-means*. In our experiments, an initialization into K clusters obtained with *k-means++* [4] was used. *k-means++* is basically identical to the *k-means* algorithm, except for the selection of initial centers. More concretely, *k-means++* starts with

allocating one cluster center at random and then searches for the next ones which will be selected with a probability proportional to the distance to the closest center already chosen. The essential idea is to make all centers that would be selected as far away as possible from each other. It should be noted that *k-means++* also uses random initialization as a starting point, so it can give different results on different runs. To overcome the issue of poor initialization, we propose to run *k-means++* several times and use the labels that yield the best compactness (the sum of squared distances from each point to their corresponding center) to initialize the $q_Z(z_j)$'s and thus update the ρ_k 's. From the initialization of $q_Z(\mathbf{z})$ and ρ , the VB-E- τ and VB-E- θ^* steps can then be derived. This simple scheme has the advantage of accelerating the convergence of the VBEM algorithm.

Acknowledgements This article was developed in the framework of the Grenoble Alpes Data Institute, supported by the French National Research Agency under the “Investissements d’avenir” program (ANR-15-IDEX-02).

References

1. Achanta, R., Shaji, A., Smith, K., Lucchi, A., Fua, P., Süsstrunk, S.: SLIC superpixels compared to state-of-the-art superpixel methods. *IEEE Trans. Pattern Anal. Mach. Intell.* **34**(11), 2274–2282 (2012)
2. Albughdadi, M., Chaîri, L., Tourneret, J., Forbes, F., Ciuciu, P.: A Bayesian non-parametric hidden Markov random model for hemodynamic brain parcellation. *Signal Processing* **135**, 132–146 (2017)
3. Arbelaez, P., Maire, M., Fowlkes, C., Malik, J.: Contour detection and hierarchical image segmentation. *IEEE Trans. Pattern Anal. Mach. Intell.* **33**(5), 898–916 (2011)
4. Arthur, D., Vassilvitskii, S.: K-means++: The advantages of careful seeding. In: *Proceedings of the Eighteenth Annual ACM-SIAM Symposium on Discrete Algorithms, SODA '07*, pp. 1027–1035. Society for Industrial and Applied Mathematics, Philadelphia, PA, USA (2007). URL <http://dl.acm.org/citation.cfm?id=1283383.1283494>
5. Beal, M., Ghahramani, Z.: The variational Bayesian EM Algorithm for incomplete data: with application to scoring graphical model structures. In: J.M. Bernardo, M.J. Bayarri, J.O. Berger, A.P. Dawid, D. Heckerman, A.F.M. Smith, M. West (eds.) *Bayesian Statistics*, pp. 453–464. Oxford University Press (2003)
6. Beal, M.J.: Variational algorithms for approximate Bayesian inference. PhD thesis (2003)
7. Besag, J.: Spatial interaction and the statistical analysis of lattice systems. *Journal of the Royal Statistical Society. Series B (Methodological)* pp. 192–236 (1974)
8. Blei, D.M., Jordan, M.I.: Variational inference for Dirichlet process mixtures. *Bayesian Anal.* **1**(1), 121–143 (2006)
9. Celeux, G., Forbes, F., Peyrard, N.: EM procedures using mean field-like approximations for Markov model-based image segmentation. *Pattern Recognition* **36**, 131–144 (2003)
10. Chaari, L., Vincent, T., Forbes, F., Dojat, M., Ciuciu, P.: Fast joint detection-estimation of evoked brain activity in event-related fmri using a variational approach. *IEEE Trans. Med. Imag.* **32**(5), 821–837 (2013)
11. Chandler, D.: Introduction to modern statistical mechanics. Oxford University Press, New York, Oxford (1987). URL <http://opac.inria.fr/record=b1081336>
12. Chatzis, S.P.: A Markov random field-regulated Pitman-Yor process prior for spatially constrained data clustering. *Pattern Recognition* **46**(6), 1595–1603 (2013)
13. Chatzis, S.P., Tschepnakis, G.: The infinite hidden Markov random field model. *IEEE Trans. Neural Networks* **21**(6), 1004–1014 (2010)
14. Corduneanu, A., Bishop, C.M.: Variational Bayesian Model Selection for Mixture Distributions. In: *Proceedings Eighth International Conference on Artificial Intelligence and Statistics*, pp. 27–34. Morgan Kaufmann (2001)
15. De Blasi, P., Favaro, S., Lijoi, A., Mena, R.H., Prünster, I., Ruggiero, M.: Are Gibbs-type priors the most natural generalization of the Dirichlet process? *IEEE Trans. Pattern Anal. Mach. Intell.* **37**(2), 212–229 (2015)
16. Favaro, S., Lijoi, A., Nava, C., Nipoti, B., Prünster, I., Teh, Y.W.: On the Stick-Breaking Representation for Homogeneous NRMIs. *Bayesian Anal.* **11**(3), 697–724 (2016)
17. Ferguson, T.S.: A Bayesian analysis of some nonparametric problems. *The Annals of Statistics* pp. 209–230 (1973)
18. Forbes, F., Peyrard, N.: Hidden Markov Random Field model selection criteria based on mean field-like approximations. *IEEE Trans. Pattern Anal. Mach. Intell.* **25**(9), 1089–1101 (2003)
19. Ghosal, S., Van der Vaart, A.: *Fundamentals of nonparametric Bayesian inference*, vol. 44. Cambridge University Press (2017)
20. Green, P.: Reversible jump Markov chain Monte Carlo computation and Bayesian model determination. *Biometrika* **82**, 711–732 (1995)
21. Ishwaran, H., James, L.F.: Gibbs sampling methods for stick-breaking priors. *Journal of the American Statistical Association* **96**(453), 161–173 (2001)
22. Johnson, T.D., Liu, Z., Bartsch, A.J., Nichols, T.E.: A Bayesian non-parametric Potts model with application to pre-surgical fMRI data. *Statistical Methods in Medical Research* **22**(4), 364–381 (2013)
23. McLachlan, G., Krishnan, T.: *The EM Algorithm and Extensions*. Wiley (1996)
24. Miller, J.W., Harrison, M.T.: Mixture models with a prior on the number of components. *Journal of the American Statistical Association* **113**, 340–356 (2018)
25. Murphy, K.P.: Conjugate bayesian analysis of the gaussian distribution. *def 1(2σ2)*, 16 (2007)
26. Neal, R.M., Hinton, G.E.: A view of the EM algorithm that justifies incremental, sparse and other variants. In: Jordan (ed.) *Lear. in Graph. Mod.*, pp. 355–368 (1998)
27. Orbanz, P., Buhmann, J.M.: Nonparametric Bayesian image segmentation. *International Journal of Computer Vision* **77**(1-3), 25–45 (2008)
28. Pitman, J.: Combinatorial stochastic processes, *Lecture Notes in Mathematics*, vol. 1875. Springer-Verlag, Berlin (2006). Lectures from the 32nd Summer School on Probability Theory held in Saint-Flour, July 7–24, 2002
29. Pitman, J., Yor, M.: The Two-Parameter Poisson-Dirichlet Distribution Derived from a Stable Subordinator. *The Annals of Probability* **25**(2), 855–900 (1997)
30. Rand, W.M.: Objective criteria for the evaluation of clustering methods. *Journal of the American Statistical Association* **66**(336), 846–850 (1971)
31. Sethuraman, J.: A constructive definition of dirichlet priors. *Statistica Sinica* **4**(2), 639–650 (1994)

32. Shyr, A., Darrell, T., Jordan, M.I., Urtasun, R.: Supervised hierarchical Pitman-Yor process for natural scene segmentation. In: proceedings of CVPR 2011, pp. 2281–2288 (2011)
33. da Silva, A.R.F.: A Dirichlet process mixture model for brain MRI tissue classification. *Medical Image Analysis* **11**(2), 169–182 (2007)
34. Sodjo, J., Giremus, A., Dobigeon, N., Giovannelli, J.F.: A generalized Swendsen-Wang algorithm for Bayesian nonparametric joint segmentation of multiple images. In: IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP), pp. 1882–1886. IEEE, La Nouvelle Orléans, LA, United States (2017)
35. Stoehr, J.: A review on statistical inference methods for discrete Markov random fields. arXiv e-prints arXiv:1704.03331 (2017)
36. Sudderth, E.B., Jordan, M.I.: Shared segmentation of natural scenes using dependent Pitman-Yor processes. In: Advances in Neural Information Processing Systems 21, Proceedings of the Twenty-Second Annual Conference on Neural Information Processing Systems, Vancouver, British Columbia, Canada, December 8–11, 2008, pp. 1585–1592 (2008)
37. Teh, Y.W.: A Bayesian interpretation of interpolated Kneser-Ney. Technical report (2006)
38. Unnikrishnan, R., Pantofaru, C., Hebert, M.: A measure for objective evaluation of image segmentation algorithms. In: 2005 IEEE Computer Society Conference on Computer Vision and Pattern Recognition (CVPR'05) - Workshops, pp. 34–34 (2005)
39. Varma, M., Zisserman, A.: A statistical approach to texture classification from single images. *International Journal of Computer Vision* **62**(1), 61–81 (2005)
40. Wallach, H., Jensen, S., Dicker, L., Heller, K.: An Alternative Prior Process for Nonparametric Bayesian Clustering. In: Y.W. Teh, M. Titterton (eds.) Proceedings of the Thirteenth International Conference on Artificial Intelligence and Statistics, *Proceedings of Machine Learning Research*, vol. 9, pp. 892–899. PMLR, Chia Laguna Resort, Sardinia, Italy (2010)
41. Wang, C., Blei, D.M.: Truncation-free stochastic variational inference for Bayesian nonparametric models. In: Proceedings of the 25th International Conference on Neural Information Processing Systems - Volume 1, NIPS'12, pp. 413–421 (2012)
42. Xu, D., Caron, F., Doucet, A.: Bayesian nonparametric image segmentation using a generalized Swendsen-Wang algorithm. ArXiv e-prints (2016)