
Analysis of Variational Bayesian Latent Dirichlet Allocation: Weaker Sparsity than MAP

Shinichi Nakajima

Berlin Big Data Center, TU Berlin
Berlin 10587 Germany
nakajima@tu-berlin.de

Issei Sato

University of Tokyo
Tokyo 113-0033 Japan
sato@r.dl.itc.u-tokyo.ac.jp

Masashi Sugiyama

University of Tokyo
Tokyo 113-0033, Japan
sugi@k.u-tokyo.ac.jp

Kazuho Watanabe

Toyohashi University of Technology
Aichi 441-8580 Japan
wkazuho@cs.tut.ac.jp

Hiroko Kobayashi

Nikon Corporation
Kanagawa 244-8533 Japan
hiroko.kobayashi@nikon.com

Abstract

Latent Dirichlet allocation (LDA) is a popular generative model of various objects such as texts and images, where an object is expressed as a mixture of latent *topics*. In this paper, we theoretically investigate *variational Bayesian* (VB) learning in LDA. More specifically, we analytically derive the leading term of the VB free energy under an asymptotic setup, and show that there exist transition thresholds in Dirichlet hyperparameters around which the sparsity-inducing behavior drastically changes. Then we further theoretically reveal the notable phenomenon that *VB tends to induce weaker sparsity than MAP* in the LDA model, which is opposed to other models. We experimentally demonstrate the practical validity of our asymptotic theory on real-world *Last.FM* music data.

1 Introduction

Latent Dirichlet allocation (LDA) [5] is a generative model successfully used in various applications such as text analysis [5], image analysis [15], genomics [6, 4], human activity analysis [12], and collaborative filtering [14, 20]¹. Given word occurrences of documents in a corpora, LDA expresses each document as a mixture of multinomial distributions, each of which is expected to capture a *topic*. The extracted topics provide bases in a low-dimensional feature space, in which each document is compactly represented. This topic expression was shown to be useful for solving various tasks including classification [15], retrieval [26], and recommendation [14].

Since rigorous Bayesian inference is computationally intractable in the LDA model, various approximation techniques such as *variational Bayesian* (VB) learning [3, 7] are used. Previous theoretical studies on VB learning revealed that VB tends to produce sparse solutions, e.g., in mixture models [24, 25, 13], hidden Markov models [11], Bayesian networks [23], and fully-observed matrix factorization [17]. Here, we mean by sparsity that VB exhibits the automatic relevance determination

¹ For simplicity, we use the terminology in text analysis below. However, the range of application of our theory given in this paper is not limited to texts.

(ARD) effect [19], which automatically prunes irrelevant degrees of freedom under non-informative or weakly sparse prior. Therefore, it is naturally expected that VB-LDA also produces a sparse solution (in terms of topics). However, it is often observed that VB-LDA does not generally give sparse solutions.

In this paper, we attempt to clarify this gap by theoretically investigating the sparsity-inducing mechanism of VB-LDA. More specifically, we first analytically derive the leading term of the VB free energy in some asymptotic limits, and show that there exist transition thresholds in Dirichlet hyperparameters around which the sparsity-inducing behavior changes drastically. We then analyze the behavior of MAP and its variants in a similar way, and show that *the VB solution is less sparse than the MAP solution* in the LDA model. This phenomenon is completely opposite to other models such as mixture models [24, 25, 13], hidden Markov models [11], Bayesian networks [23], and fully-observed matrix factorization [17], where VB tends to induce stronger sparsity than MAP. We numerically demonstrate the practical validity of our asymptotic theory using artificial and real-world *Last.FM* music data for collaborative filtering, and further discuss the peculiarity of the LDA model in terms of sparsity.

The free energy of VB-LDA was previously analyzed in [16], which evaluated the advantage of collapsed VB [21] over the original VB learning. However, that work focused on the difference between VB and collapsed VB, and neither the absolute free energy nor the sparsity was investigated. The update rules of VB was compared with those of MAP [2]. However, that work is based on approximation, and rigorous analysis was not made. To the best of our knowledge, our paper is the first work that theoretically elucidates the sparsity-inducing mechanism of VB-LDA.

2 Formulation

In this section, we introduce the latent Dirichlet allocation model and variational Bayesian learning.

2.1 Latent Dirichlet Allocation

Suppose that we observe M documents, each of which consists of $N^{(m)}$ words. Each word is included in a vocabulary with size L . We assume that each word is associated with one of the H topics, which is not observed. We express the word occurrence by an L -dimensional indicator vector $\mathbf{w}$, where one of the entries is equal to one and the others are equal to zero. Similarly, we express the topic occurrence as an H -dimensional indicator vector $\mathbf{z}$. We define the following functions that give the item numbers chosen by $\mathbf{w}$ and $\mathbf{z}$, respectively:

$$\dot{l}(\mathbf{w}) = l \text{ if } w_l = 1 \text{ and } w_{l'} = 0 \text{ for } l' \neq l, \quad \dot{h}(\mathbf{z}) = h \text{ if } z_h = 1 \text{ and } z_{h'} = 0 \text{ for } h' \neq h.$$

In the latent Dirichlet allocation (LDA) model [5], the word occurrence $\mathbf{w}^{(n,m)}$ of the n -th position in the m -th document is assumed to follow the multinomial distribution:

$$p(\mathbf{w}^{(n,m)} | \boldsymbol{\Theta}, \mathbf{B}) = \prod_{l=1}^L \left((\mathbf{B}\boldsymbol{\Theta}^\top)_{l,m} \right)^{w_l^{(n,m)}} = (\mathbf{B}\boldsymbol{\Theta}^\top)_{\dot{l}(\mathbf{w}^{(n,m)}),m}, \quad (1)$$

where $\boldsymbol{\Theta} \in [0, 1]^{M \times H}$ and $\mathbf{B} \in [0, 1]^{L \times H}$ are parameter matrices to be estimated. The rows of $\boldsymbol{\Theta}$ and the columns of $\mathbf{B}$ are probability mass vectors that sum up to one. We denote a column vector of a matrix by a bold lowercase letter, and a row vector by a bold lowercase letter with a tilde, i.e.,

$$\boldsymbol{\Theta} = (\boldsymbol{\theta}_1, \dots, \boldsymbol{\theta}_H) = (\tilde{\boldsymbol{\theta}}_1, \dots, \tilde{\boldsymbol{\theta}}_M)^\top, \quad \mathbf{B} = (\boldsymbol{\beta}_1, \dots, \boldsymbol{\beta}_H) = (\tilde{\boldsymbol{\beta}}_1, \dots, \tilde{\boldsymbol{\beta}}_L)^\top.$$

With this notation, $\tilde{\boldsymbol{\theta}}_m$ denotes the topic distribution of the m -th document, and $\boldsymbol{\beta}_h$ denotes the word distribution of the h -th topic.

Given the topic occurrence latent variable $\mathbf{z}^{(n,m)}$, the complete likelihood is written as

$$p(\mathbf{w}^{(n,m)}, \mathbf{z}^{(n,m)} | \boldsymbol{\Theta}, \mathbf{B}) = p(\mathbf{w}^{(n,m)} | \mathbf{z}^{(n,m)}, \mathbf{B}) p(\mathbf{z}^{(n,m)} | \boldsymbol{\Theta}), \quad (2)$$

where $p(\mathbf{w}^{(n,m)} | \mathbf{z}^{(n,m)}, \mathbf{B}) = \prod_{l=1}^L \prod_{h=1}^H (B_{l,h})^{w_l^{(n,m)} z_h^{(n,m)}}$, $p(\mathbf{z}^{(n,m)} | \boldsymbol{\Theta}) = \prod_{h=1}^H (\boldsymbol{\theta}_{m,h})^{z_h^{(n,m)}}$.

We assume the Dirichlet prior on $\boldsymbol{\Theta}$ and $\mathbf{B}$:

$$p(\boldsymbol{\Theta} | \alpha) \propto \prod_{m=1}^M \prod_{h=1}^H (\boldsymbol{\theta}_{m,h})^{\alpha-1}, \quad p(\mathbf{B} | \eta) \propto \prod_{h=1}^H \prod_{l=1}^L (B_{l,h})^{\eta-1}, \quad (3)$$

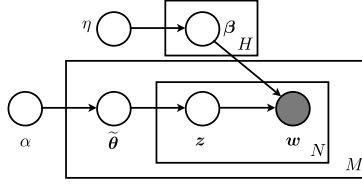


Figure 1: Graphical model of LDA.

where α and η are hyperparameters that control the prior sparsity. We can make α dependent on m and/or h , and η dependent on l and/or h , and they can be estimated from observation. However, we fix those hyperparameters as given constants for simplicity in our analysis below. Figure 1 shows the graphical model of LDA.

2.2 Variational Bayesian Learning

The Bayes posterior of LDA is written as

$$p(\Theta, B, \{z^{(n,m)}\} | \{w^{(n,m)}\}, \alpha, \eta) = \frac{p(\{w^{(n,m)}\}, \{z^{(n,m)}\} | \Theta, B) p(\Theta | \alpha) p(B | \eta)}{p(\{w^{(n,m)}\})}, \quad (4)$$

where $p(\{w^{(n,m)}\}) = \int p(\{w^{(n,m)}\}, \{z^{(n,m)}\} | \Theta, B) p(\Theta | \alpha) p(B | \eta) d\Theta dB d\{z^{(n,m)}\}$ is intractable to compute and thus requires some approximation method. In this paper, we focus on the variational Bayesian (VB) approximation and investigate its behavior theoretically.

In the VB approximation, we assume that our approximate posterior is factorized as

$$q(\Theta, B, \{z^{(n,m)}\}) = q(\Theta, B) q(\{z^{(n,m)}\}), \quad (5)$$

and minimize the free energy:

$$F = \left\langle \log \frac{q(\Theta, B, \{z^{(n,m)}\})}{p(\{w^{(n,m)}\}, \{z^{(n,m)}\} | \Theta, B) p(\Theta | \alpha) p(B | \eta)} \right\rangle_{q(\Theta, B, \{z^{(n,m)}\})}, \quad (6)$$

where $\langle \cdot \rangle_p$ denotes the expectation over the distribution p . This amounts to finding the distribution that is closest to the Bayes posterior (4) under the constraint (5). Using the variational method, we can obtain the following stationary condition:

$$q(\Theta) \propto p(\Theta | \alpha) \exp \left\langle \log p(\{w^{(n,m)}\}, \{z^{(n,m)}\} | \Theta, B) \right\rangle_{q(B) q(\{z^{(n,m)}\})}, \quad (7)$$

$$q(B) \propto p(B | \eta) \exp \left\langle \log p(\{w^{(n,m)}\}, \{z^{(n,m)}\} | \Theta, B) \right\rangle_{q(\Theta) q(\{z^{(n,m)}\})}, \quad (8)$$

$$q(\{z^{(n,m)}\}) \propto \exp \left\langle \log p(\{w^{(n,m)}\}, \{z^{(n,m)}\} | \Theta, B) \right\rangle_{q(\Theta) q(B)}. \quad (9)$$

From this, we can confirm that $\{q(\tilde{\theta}_m)\}$ and $\{q(\beta_h)\}$ follow the Dirichlet distribution and $\{q(z^{(n,m)})\}$ follows the multinomial distribution:

$$q(\Theta) \propto \prod_{m=1}^M \prod_{h=1}^H (\Theta_{m,h})^{\check{\Theta}_{m,h}-1}, \quad q(B) \propto \prod_{h=1}^H \prod_{l=1}^L (B_{l,h})^{\check{B}_{l,h}-1}, \quad (10)$$

$$q(\{z^{(n,m)}\}) = \prod_{m=1}^M \prod_{n=1}^{N^{(m)}} \prod_{h=1}^H (\hat{z}_h^{(n,m)})^{z_h^{(n,m)}}, \quad (11)$$

where, for $\psi(\cdot)$ denoting the Digamma function, the variational parameters satisfy

$$\check{\Theta}_{m,h} = \alpha + \sum_{n=1}^{N^{(m)}} \hat{z}_h^{(n,m)}, \quad \check{B}_{l,h} = \eta + \sum_{m=1}^M \sum_{n=1}^{N^{(m)}} w_l^{(n,m)} \hat{z}_h^{(n,m)}, \quad (12)$$

$$\hat{z}_h^{(n,m)} = \frac{\exp \left\{ \Psi(\check{\Theta}_{m,h}) + \sum_{l=1}^L w_l^{(n,m)} (\Psi(\check{B}_{l,h}) - \Psi(\sum_{l'=1}^L \check{B}_{l',h})) \right\}}{\sum_{h'=1}^H \exp \left\{ \Psi(\check{\Theta}_{m,h'}) + \sum_{l=1}^L w_l^{(n,m)} (\Psi(\check{B}_{l,h'}) - \Psi(\sum_{l'=1}^L \check{B}_{l',h'})) \right\}}. \quad (13)$$

2.3 Partially Bayesian Learning and MAP Estimation

We can partially apply VB learning by approximating the posterior of Θ or B by the delta function. This approach is called the *partially Bayesian* (PA) learning [18], whose behavior was analyzed

and compared with VB in fully-observed matrix factorization. We call it PBA learning if Θ is marginalized and B is point-estimated, and PBB learning if B is marginalized and Θ is point-estimated. Note that the original VB algorithm for LDA proposed by [5] corresponds to PBA in our terminology. We also analyze the behavior of MAP estimation, where both of Θ and B are point-estimated. This corresponds to the *probabilistic latent semantic analysis* (pLSA) model [10], if we assume the flat prior $\alpha = \eta = 1$ [8].

3 Theoretical Analysis

In this section, we first give an explicit form of the free energy in the LDA model. We then investigate its asymptotic behavior for VB learning, and further conduct similar analyses to the PBA, PBB, and MAP methods. Finally, we discuss the sparsity-inducing mechanism of these learning methods, and the relation to previous theoretical studies.

3.1 Explicit Form of Free Energy

We first express the free energy (6) as a function of the variational parameters $\check{\Theta}$ and $\check{B}$:

$$F = R + Q, \quad \text{where} \quad (14)$$

$$\begin{aligned} R &= \left\langle \log \frac{q(\Theta)q(B)}{p(\Theta|\alpha)p(B|\eta)} \right\rangle_{q(\Theta, B)} \\ &= \sum_{m=1}^M \left(\log \frac{\Gamma(\sum_{h=1}^H \check{\Theta}_{m,h})}{\prod_{h=1}^H \Gamma(\check{\Theta}_{m,h})} \frac{\Gamma(\alpha)^H}{\Gamma(H\alpha)} + \sum_{h=1}^H \left(\check{\Theta}_{m,h} - \alpha \right) \left(\Psi(\check{\Theta}_{m,h}) - \Psi(\sum_{h'=1}^H \check{\Theta}_{m,h'}) \right) \right) \\ &\quad + \sum_{h=1}^H \left(\log \frac{\Gamma(\sum_{l=1}^L \check{B}_{l,h})}{\prod_{l=1}^L \Gamma(\check{B}_{l,h})} \frac{\Gamma(\eta)^L}{\Gamma(L\eta)} + \sum_{l=1}^L \left(\check{B}_{l,h} - \eta \right) \left(\Psi(\check{B}_{l,h}) - \Psi(\sum_{l'=1}^L \check{B}_{l',h}) \right) \right), \end{aligned} \quad (15)$$

$$\begin{aligned} Q &= \left\langle \log \frac{q(\{z^{(n,m)}\})}{p(\{w^{(n,m)}\}, \{z^{(n,m)}\}|\Theta, B)} \right\rangle_{q(\Theta, B, \{z^{(n,m)}\})} \\ &= - \sum_{m=1}^M N^{(m)} \sum_{l=1}^L V_{l,m} \log \left(\sum_{h=1}^H \frac{\exp(\Psi(\check{\Theta}_{m,h}))}{\exp(\Psi(\sum_{h'=1}^H \check{\Theta}_{m,h'}))} \frac{\exp(\Psi(\check{B}_{l,h}))}{\exp(\Psi(\sum_{l'=1}^L \check{B}_{l',h}))} \right). \end{aligned} \quad (16)$$

Here, $V \in \mathbb{R}^{L \times M}$ is the empirical word distribution matrix with its entries given by $V_{l,m} = \frac{1}{N^{(m)}} \sum_{n=1}^{N^{(m)}} w_l^{(n,m)}$. Note that we have eliminated the variational parameters $\{\hat{z}^{(n,m)}\}$ for the topic occurrence latent variables by using the stationary condition (13).

3.2 Asymptotic Analysis of VB Solution

Below, we investigate the leading term of the free energy in the asymptotic limit when $N \equiv \min_m N^{(m)} \rightarrow \infty$. Unlike the previous analysis for latent variable models [24], we do not assume $L, M \ll N$, but $1 \ll L, M, N$ at this point. This amounts to considering the asymptotic limit when $L, M, N \rightarrow \infty$ with a fixed mutual ratio, or equivalently, assuming $L, M \sim O(N)$. Throughout the paper, H is set at $H = \min(L, M)$ (i.e., the matrix $B\Theta^\top$ can express any multinomial distribution). We assume that the word distribution matrix V is a sample from the multinomial distribution with the *true* parameter $U^* \in \mathbb{R}^{L \times M}$ whose rank is $H^* \sim O(1)$, i.e., $U^* = B^* \Theta^{*\top}$ where $\Theta^* \in \mathbb{R}^{M \times H^*}$ and $B^* \in \mathbb{R}^{L \times H^*}$.² We assume that $\alpha, \eta \sim O(1)$.

The stationary condition (12) leads to the following lemma (the proof is given in Appendix A):

Lemma 1 Let $\hat{B}\hat{\Theta}^\top = \langle B\Theta^\top \rangle_{q(\Theta, B)}$. Then, it holds that

$$\langle (B\Theta^\top - \hat{B}\hat{\Theta}^\top)_{l,m}^2 \rangle_{q(\Theta, B)} = O_p(N^{-2}), \quad (17)$$

$$Q = - \sum_{m=1}^M N^{(m)} \sum_{l=1}^L V_{l,m} \log(\hat{B}\hat{\Theta}^\top)_{l,m} + O_p(M), \quad (18)$$

where $O_p(\cdot)$ denotes the order in probability.

² More precisely, $U^* = B^* \Theta^{*\top} + O(N^{-1})$ is sufficient.

Eq.(17) implies the convergence of the posterior. Let

$$\hat{J} = \sum_{l=1}^L \sum_{m=1}^M \kappa \left((\hat{\mathbf{B}}\hat{\boldsymbol{\Theta}}^\top)_{l,m} \neq (\mathbf{B}^*\boldsymbol{\Theta}^{*\top})_{l,m} + O_p(N^{-1}) \right) \quad (19)$$

be the number of entries of $\hat{\mathbf{B}}\hat{\boldsymbol{\Theta}}^\top$ that does not converge to the true value. Here, we denote by $\kappa(\cdot)$ the indicator function equal to one if the *event* is true, and zero otherwise. Then, Eq.(18) leads to the following lemma:

Lemma 2 *Q is minimized when $\hat{\mathbf{B}}\hat{\boldsymbol{\Theta}}^\top = \mathbf{B}^*\boldsymbol{\Theta}^{*\top} + O_p(N^{-1})$, and it holds that*

$$Q = S + O_p(\hat{J}N + M), \quad \text{where} \\ S = -\log p(\{\mathbf{w}^{(n,m)}\}, \{\mathbf{z}^{(n,m)}\} | \boldsymbol{\Theta}^*, \mathbf{B}^*) = -\sum_{m=1}^M N^{(m)} \sum_{l=1}^L V_{l,m} \log(\mathbf{B}^*\boldsymbol{\Theta}^*)_{l,m}.$$

Lemma 2 simply states that Q/N converges to the normalized entropy S/N of the true distribution (which is the lowest achievable value with probability 1), if and only if VB converges to the true distribution (i.e., $\hat{J} = 0$).

Let $\hat{H} = \sum_{h=1}^H \kappa(\frac{1}{M} \sum_{m=1}^M \hat{\Theta}_{m,h} \sim O_p(1))$ be the number of topics used in the whole corpus, $\hat{M}^{(h)} = \sum_{m=1}^M \kappa(\hat{\Theta}_{m,h} \sim O_p(1))$ be the number of documents that contain the h -th topic, and $\hat{L}^{(h)} = \sum_{l=1}^L \kappa(\hat{B}_{l,h} \sim O_p(1))$ be the number of words of which the h -th topic consist. We have the following lemma (the proof is given in Appendix B):

Lemma 3 *R is written as follows:*

$$R = \left\{ M \left(H\alpha - \frac{1}{2} \right) + \hat{H} \left(L\eta - \frac{1}{2} \right) - \sum_{h=1}^{\hat{H}} \left(\hat{M}^{(h)} \left(\alpha - \frac{1}{2} \right) + \hat{L}^{(h)} \left(\eta - \frac{1}{2} \right) \right) \right\} \log N \\ + (H - \hat{H}) \left(L\eta - \frac{1}{2} \right) \log L + O_p(H(M + L)). \quad (20)$$

Since we assumed that the true matrices $\boldsymbol{\Theta}^*$ and $\mathbf{B}^*$ are of the rank of H^* , $\hat{H} = H^* \sim O(1)$ is sufficient for the VB posterior to converge to the *true* distribution. However, $\hat{H}$ can be much larger than H^* with $\langle \mathbf{B}\boldsymbol{\Theta}^\top \rangle_{q(\boldsymbol{\Theta}, \mathbf{B})}$ unchanged because of the non-identifiability of matrix factorization—duplicating topics with divided weights, for example, does not change the distribution.

Based on Lemma 2 and Lemma 3, we obtain the following theorem (the proof is given in Appendix C):

Theorem 1 *In the limit when $N \rightarrow \infty$ with $L, M \sim O(1)$, it holds that $\hat{J} = 0$ with probability 1, and*

$$F = S + \left\{ M \left(H\alpha - \frac{1}{2} \right) + \hat{H} \left(L\eta - \frac{1}{2} \right) - \sum_{h=1}^{\hat{H}} \left(\hat{M}^{(h)} \left(\alpha - \frac{1}{2} \right) + \hat{L}^{(h)} \left(\eta - \frac{1}{2} \right) \right) \right\} \log N \\ + O_p(1).$$

In the limit when $N, M \rightarrow \infty$ with $\frac{M}{N}, L \sim O(1)$, it holds that $\hat{J} = o_p(\log N)$, and

$$F = S + \left\{ M \left(H\alpha - \frac{1}{2} \right) - \sum_{h=1}^{\hat{H}} \hat{M}^{(h)} \left(\alpha - \frac{1}{2} \right) \right\} \log N + o_p(N \log N).$$

In the limit when $N, L \rightarrow \infty$ with $\frac{L}{N}, M \sim O(1)$, it holds that $\hat{J} = o_p(\log N)$, and

$$F = S + HL\eta \log N + o_p(N \log N).$$

In the limit when $N, L, M \rightarrow \infty$ with $\frac{L}{N}, \frac{M}{N} \sim O(1)$, it holds that $\hat{J} = o_p(N \log N)$, and

$$F = S + H(M\alpha + L\eta) \log N + o_p(N^2 \log N).$$

Since Eq.(17) was shown to hold, the predictive distribution converges to the true distribution if $\hat{J} = 0$. Accordingly, Theorem 1 states that the consistency holds in the limit when $N \rightarrow \infty$ with $L, M \sim O(1)$.

Theorem 1 also implies that, in the asymptotic limits with small $L \sim O(1)$, the leading term depends on $\hat{H}$, meaning that it dominates the topic sparsity of the VB solution. We have the following corollary (the proof is given in Appendix D):

Table 1: Sparsity thresholds of VB, PBA, PBB, and MAP methods (see Theorem 2). The first four columns show the thresholds $(\underline{\alpha}_{\text{sparse}}, \underline{\alpha}_{\text{dense}})$, of which the function forms depend on the range of η , in the limit when $N \rightarrow \infty$ with $L, M \sim O(1)$. A single value is shown if $\underline{\alpha}_{\text{sparse}} = \underline{\alpha}_{\text{dense}}$. The last column shows the threshold $\underline{\alpha}_{M \rightarrow \infty}$ in the limit when $N, M \rightarrow \infty$ with $\frac{M}{N}, L \sim O(1)$.

η range	$(\underline{\alpha}_{\text{sparse}}, \underline{\alpha}_{\text{dense}})$				$\underline{\alpha}_{M \rightarrow \infty}$
	$0 < \eta \leq \frac{1}{2L}$	$\frac{1}{2L} < \eta \leq \frac{1}{2}$	$\frac{1}{2} < \eta < 1$	$1 \leq \eta < \infty$	$0 < \eta < \infty$
VB	$\frac{1}{2} - \frac{\frac{1}{2} - L\eta}{\min_h M^{*(h)}}$	$\frac{1}{2} + \frac{L\eta - \frac{1}{2}}{\max_h M^{*(h)}}$	$\left(\frac{1}{2} + \frac{L-1}{2 \max_h M^{*(h)}}, \frac{1}{2} + \frac{L\eta - \frac{1}{2}}{\min_h M^{*(h)}}\right)$		$\frac{1}{2}$
PBA	—			$\left(\frac{1}{2}, \frac{1}{2} + \frac{L(\eta-1)}{\min_h M^{*(h)}}\right)$	$\frac{1}{2}$
PBB	1	$1 + \frac{L\eta - \frac{1}{2}}{\max_h M^{*(h)}}$	$\left(1 + \frac{L-1}{2 \max_h M^{*(h)}}, 1 + \frac{L\eta - \frac{1}{2}}{\min_h M^{*(h)}}\right)$		1
MAP	—			$\left(1, 1 + \frac{L(\eta-1)}{\min_h M^{*(h)}}\right)$	1

Corollary 1 Let $M^{*(h)} = \sum_{m=1}^M \kappa(\Theta_{m,h}^* \sim O(1))$ and $L^{*(h)} = \sum_{l=1}^L \kappa(B_{l,h}^* \sim O(1))$. Consider the limit when $N \rightarrow \infty$ with $L, M \sim O(1)$. When $0 < \eta \leq \frac{1}{2L}$, the VB solution is sparse if $\alpha < \frac{1}{2} - \frac{\frac{1}{2} - L\eta}{\min_h M^{*(h)}}$, and dense if $\alpha > \frac{1}{2} - \frac{\frac{1}{2} - L\eta}{\min_h M^{*(h)}}$. When $\frac{1}{2L} < \eta \leq \frac{1}{2}$, the VB solution is sparse if $\alpha < \frac{1}{2} + \frac{L\eta - \frac{1}{2}}{\max_h M^{*(h)}}$, and dense if $\alpha > \frac{1}{2} + \frac{L\eta - \frac{1}{2}}{\max_h M^{*(h)}}$. When $\eta > \frac{1}{2}$, the VB solution is sparse if $\alpha < \frac{1}{2} + \frac{L-1}{2 \max_h M^{*(h)}}$, and dense if $\alpha > \frac{1}{2} + \frac{L\eta - \frac{1}{2}}{\min_h M^{*(h)}}$. In the limit when $N, M \rightarrow \infty$ with $\frac{M}{N}, L \sim O(1)$, the VB solution is sparse if $\alpha < \frac{1}{2}$, and dense if $\alpha > \frac{1}{2}$.

In the case when $L, M \ll N$ and in the case when $L \ll M, N$, Corollary 1 provides information on the sparsity of the VB solution, which will be compared with other methods in Section 3.3. On the other hand, although we have successfully derived the leading term of the free energy also in the case when $M \ll L, N$ and in the case when $1 \ll L, M, N$, it unfortunately provides no information on sparsity of the solution.

3.3 Asymptotic Analysis of PBA, PBB, and MAP

By applying similar analysis to PBA learning, PBB learning, and MAP estimation, we can obtain the following theorem (the proof is given in Appendix E):

Theorem 2 In the limit when $N \rightarrow \infty$ with $L, M \sim O(1)$, the solution is sparse if $\alpha < \underline{\alpha}_{\text{sparse}}$, and dense if $\alpha > \underline{\alpha}_{\text{dense}}$. In the limit when $N, M \rightarrow \infty$ with $\frac{M}{N}, L \sim O(1)$, the solution is sparse if $\alpha < \underline{\alpha}_{M \rightarrow \infty}$, and dense if $\alpha > \underline{\alpha}_{M \rightarrow \infty}$. Here, $\underline{\alpha}_{\text{sparse}}$, $\underline{\alpha}_{\text{dense}}$, and $\underline{\alpha}_{M \rightarrow \infty}$ are given in Table 1.

A notable finding from Table 1 is that the threshold that determines the topic sparsity of PBB-LDA is (most of the case exactly) $\frac{1}{2}$ larger than the threshold of VB-LDA. The same relation is observed between MAP-LDA and PBA-LDA. From these, we can conclude that point-estimating Θ , instead of integrating it out, increases the threshold by $\frac{1}{2}$ in the LDA model. We will validate this observation by numerical experiments in Section 4.

3.4 Discussion

The above theoretical analysis (Theorem 2) showed that VB tends to induce weaker sparsity than MAP in the LDA model³, i.e., VB requires sparser prior (smaller α) than MAP to give a sparse solution (mean of the posterior). This phenomenon is completely opposite to other models such as mixture models [24, 25, 13], hidden Markov models [11], Bayesian networks [23], and fully-observed matrix factorization [17], where VB tends to induce stronger sparsity than MAP. This phenomenon might be partly explained as follows: In the case of mixture models, the sparsity threshold depends on the degree of freedom of a single component [24]. This is reasonable because

³ Although this tendency was previously pointed out [2] by using the approximation $\exp(\psi(n)) \approx n - \frac{1}{2}$ and comparing the stationary condition, our result has first clarified the sparsity behavior of the solution based on the asymptotic free energy analysis without using such an approximation.

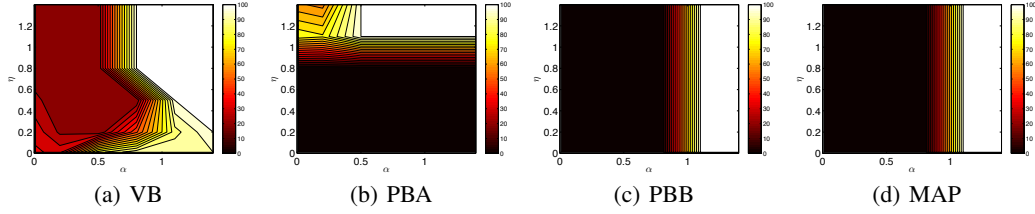


Figure 2: Estimated number $\hat{H}$ of topics by (a) VB, (b) PBA, (c) PBB, and (d) MAP, for the *artificial* data with $L = 100$, $M = 100$, $H^* = 20$, and $N \sim 10000$.

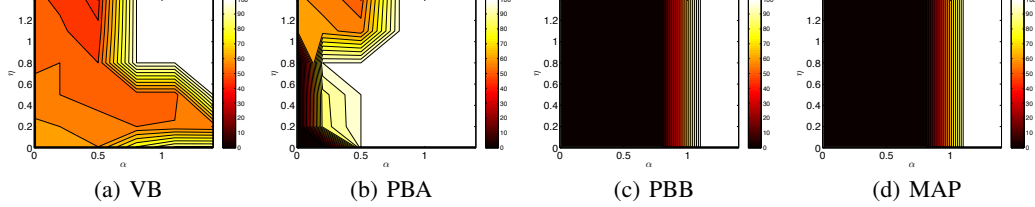


Figure 3: Estimated number $\hat{H}$ of topics for the *Last.FM* data with $L = 100$, $M = 100$, and $N \sim 700$.

adding a single component increases the model complexity by this amount. Also, in the case of LDA, adding a single topic requires additional $L + 1$ parameters. However, the added topic is shared over M documents, which could discount the increased model complexity relative to the increased data fidelity. Corollary 1, which implies the dependency of the threshold for α on L and M , might support this conjecture. However, the same applies to the matrix factorization, where VB was shown to give a sparser solution than MAP [17]. Investigation on related models, e.g., Poisson MF [9], would help us fully explain this phenomenon.

Technically, our theoretical analysis is based on the previous asymptotic studies on VB learning conducted for latent variable models [24, 25, 13, 11, 23]. However, our analysis is not just a straightforward extension of those works to the LDA model. For example, the previous analysis either implicitly [24] or explicitly [13] assumed the consistency of VB learning, while we also analyzed the consistency of VB-LDA, and showed that the consistency does not always hold (see Theorem 1). Moreover, we derived a general form of the asymptotic free energy, which can be applied to different asymptotic limits. Specifically, the standard asymptotic theory requires a large number N of words per document, compared to the number M of documents and the vocabulary size L . This may be reasonable in some collaborative filtering data such as the *Last.FM* data used in our experiments in Section 4. However, L and/or M would be comparable to or larger than N in standard text analysis.

Our general form of the asymptotic free energy also allowed us to elucidate the behavior of the VB free energy when L and/or M diverges with the same order as N . This attempt successfully revealed the sparsity of the solution for the case when M diverges while $L \sim O(1)$. However, when L diverges, we found that the leading term of the free energy does not contain interesting insight into the sparsity of the solution. Higher-order asymptotic analysis will be necessary to further understand the sparsity-inducing mechanism of the LDA model with large vocabulary.

4 Numerical Illustration

In this section, we conduct numerical experiments on artificial and real data for collaborative filtering.

The *artificial* data were created as follows. We first sample the *true* document matrix Θ^* of size $M \times H^*$ and the *true* topic matrix B^* of size $L \times H^*$. We assume that each row $\tilde{\theta}_m^*$ of Θ^* follows the Dirichlet distribution with $\alpha^* = 1/H^*$, while each column β_h^* of B^* follows the Dirichlet distribution with $\eta^* = 1/L$. The document length $N^{(m)}$ is sampled from the Poisson distribution with its mean N . The word histogram $N^{(m)}\mathbf{v}_m$ for each document is sampled from the multinomial

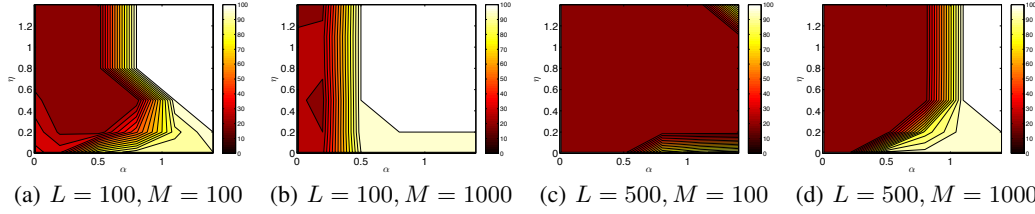


Figure 4: Estimated number $\hat{H}$ of topics by VB-LDA for the *artificial* data with $H^* = 20$ and $N \sim 10000$. For the case when $L = 500, M = 1000$, the maximum estimated rank is limited to 100 for computational reason.

distribution with the parameter specified by the m -th row vector of $\mathbf{B}^* \boldsymbol{\Theta}^{*\top}$. Thus, we obtain the $L \times M$ matrix $\mathbf{V}$, which corresponds to the empirical word distribution over M documents.

As a real-world dataset, we used the *Last.FM* dataset.⁴ *Last.FM* is a well-known social music web site, and the dataset includes the triple (“user,” “artist,” “Freq”) which was collected from the play-lists of users in the community by using a plug-in in users’ media players. This triple means that “user” played “artist” music “Freq” times, which indicates users’ preferred artists. A user and a played artist are analogous to a document and a word, respectively. We randomly chose L artists from the top 1000 frequent artists, and M users who live in the United States. To find a better local solution (which hopefully is close to the global solution), we adopted a split and merge strategy [22], and chose the local solution giving the lowest free energy among different initialization schemes.

Figure 2 shows the estimated number $\hat{H}$ of topics by different approximation methods, i.e., VB, PBA, PBB, and MAP, for the *Artificial* data with $L = 100, M = 100, H^* = 20$, and $N \sim 10000$. We can clearly see that the sparsity threshold in PBB and MAP, where $\boldsymbol{\Theta}$ is point-estimated, is larger than that in VB and PBA, where $\boldsymbol{\Theta}$ is marginalized. This result supports the statement by Theorem 2. Figure 3 shows results on the *Last.FM* data with $L = 100, M = 100$ and $N \sim 700$. We see a similar tendency to Figure 2 except the region where $\eta < 1$ for PBA, in which our theory does not predict the estimated number of topics.

Finally, we investigate how different asymptotic settings affect the topic sparsity. Figure 4 shows the sparsity dependence on L and M for the *artificial* data. The graphs correspond to the four cases mentioned in Theorem 1, i.e., (a) $L, M \ll N$, (b) $L \ll N, M$, (c) $M \ll N, L$, and (d) $1 \ll N, L, M$. Corollary 1 explains the behavior in (a) and (b), and further analysis is required to explain the behavior in (c) and (d).

5 Conclusion

In this paper, we considered variational Bayesian (VB) learning in the latent Dirichlet allocation (LDA) model and analytically derived the leading term of the asymptotic free energy. When the vocabulary size is small, our result theoretically explains the phase-transition phenomenon. On the other hand, when vocabulary size is as large as the number of words per document, the leading term tells nothing about sparsity. We need more accurate analysis to clarify the sparsity in such cases.

Throughout the paper, we assumed that the hyperparameters α and η are pre-fixed. However, α would often be estimated for each topic h , which is one of the advantages of using the LDA model in practice [5]. In the future work, we will extend the current line of analysis to the *empirical Bayesian* setting where the hyperparameters are also learned, and further elucidate the behavior of the LDA model.

Acknowledgments

The authors thank the reviewers for helpful comments. Shinichi Nakajima thanks the support from Nikon Corporation, MEXT Kakenhi 23120004, and the Berlin Big Data Center project (FKZ 01IS14013A). Masashi Sugiyama thanks the support from the JST CREST program. Kazuho Watanabe thanks the support from JSPS Kakenhi 23700175 and 25120014.

⁴<http://mtg.upf.edu/node/1671>

References

- [1] H. Alzer. On some inequalities for the Gamma and Psi functions. *Mathematics of Computation*, 66(217):373–389, 1997.
- [2] A. Asuncion, M. Welling, P. Smyth, and Y. W. Teh. On smoothing and inference for topic models. In *Proc. of UAI*, pages 27–34, 2009.
- [3] H. Attias. Inferring parameters and structure of latent variable models by variational Bayes. In *Proc. of UAI*, pages 21–30, 1999.
- [4] M. Bicego, P. Lovato, A. Ferrarini, and M. Delledonne. Biclustering of expression microarray data with topic models. In *Proc. of ICPR*, pages 2728–2731, 2010.
- [5] D. M. Blei, A. Y. Ng, and M. I. Jordan. Latent Dirichlet allocation. *Journal of Machine Learning Research*, 3:993–1022, 2003.
- [6] X. Chen, X. Hu, X. Shen, and G. Rosen. Probabilistic topic modeling for genomic data interpretation. In *2010 IEEE International Conference on Bioinformatics and Biomedicine (BIBM)*, pages 149–152, 2010.
- [7] Z. Ghahramani and M. J. Beal. Graphical models and variational methods. In *Advanced Mean Field Methods*, pages 161–177. MIT Press, 2001.
- [8] M. Girolami and A. Kaban. On an equivalence between PLSI and LDA. In *Proc. of SIGIR*, pages 433–434, 2003.
- [9] P. Gopalan, J. M. Hofman, and D. M. Blei. Scalable recommendation with Poisson factorization. *arXiv:1311.1704 [cs.LG]*, 2013.
- [10] T. Hofmann. Unsupervised learning by probabilistic latent semantic analysis. *Machine Learning*, 42:177–196, 2001.
- [11] T. Hosino, K. Watanabe, and S. Watanabe. Stochastic complexity of hidden markov models on the variational Bayesian learning. *IEICE Trans. on Information and Systems*, J89-D(6):1279–1287, 2006.
- [12] T. Huynh, Mario F., and B. Schiele. Discovery of activity patterns using topic models. In *International Conference on Ubiquitous Computing (UbiComp)*, 2008.
- [13] D. Kaji, K. Watanabe, and S. Watanabe. Phase transition of variational Bayes learning in Bernoulli mixture. *Australian Journal of Intelligent Information Processing Systems*, 11(4):35–40, 2010.
- [14] R. Krestel, P. Fankhauser, and W. Nejdl. Latent dirichlet allocation for tag recommendation. In *Proceedings of the Third ACM Conference on Recommender Systems*, pages 61–68, 2009.
- [15] F.-F. Li and P. Perona. A bayesian hierarchical model for learning natural scene categories. In *Proc. of CVPR*, pages 524–531, 2005.
- [16] I. Mukherjee and D. M. Blei. Relative performance guarantees for approximate inference in latent Dirichlet allocation. In *Advances in NIPS*, 2008.
- [17] S. Nakajima and M. Sugiyama. Theoretical analysis of Bayesian matrix factorization. *Journal of Machine Learning Research*, 12:2579–2644, 2011.
- [18] S. Nakajima, M. Sugiyama, and S. D. Babacan. On Bayesian PCA: Automatic dimensionality selection and analytic solution. In *Proc. of ICML*, pages 497–504, 2011.
- [19] R. M. Neal. *Bayesian Learning for Neural Networks*. Springer, 1996.
- [20] S. Purushotham, Y. Liu, and C. C. J. Kuo. Collaborative topic regression with social matrix factorization for recommendation systems. In *Proc. of ICML*, 2012.
- [21] Y. W. Teh, D. Newman, and M. Welling. A collapsed variational Bayesian inference algorithm for latent Dirichlet allocation. In *Advances in NIPS*, 2007.
- [22] N. Ueda, R. Nakano, Z. Ghahramani, and G. E. Hinton. SMEM algorithm for mixture models. *Neural Computation*, 12(9):2109–2128, 2000.
- [23] K. Watanabe, M. Shiga, and S. Watanabe. Upper bound for variational free energy of Bayesian networks. *Machine Learning*, 75(2):199–215, 2009.
- [24] K. Watanabe and S. Watanabe. Stochastic complexities of Gaussian mixtures in variational Bayesian approximation. *Journal of Machine Learning Research*, 7:625–644, 2006.
- [25] K. Watanabe and S. Watanabe. Stochastic complexities of general mixture models in variational Bayesian learning. *Neural Networks*, 20(2):210–219, 2007.
- [26] X. Wei and W. B. Croft. LDA-based document models for ad-hoc retrieval. In *Proc. of SIGIR*, pages 178–185, 2006.

A Proof of Lemma 1

The following bounds are known [1]:

$$\left(y - \frac{1}{2}\right) \log y - y + \frac{1}{2} \log(2\pi) \leq \log \Gamma(y) \leq \left(y - \frac{1}{2}\right) \log y - y + \frac{1}{2} \log(2\pi) + \frac{1}{12y}, \quad (21)$$

$$\log y - \frac{1}{y} \leq \Psi(y) \leq \log y - \frac{1}{2y}. \quad (22)$$

For the Dirichlet distribution $p(\boldsymbol{\theta}|\check{\boldsymbol{\theta}}) \propto \prod_{h=1}^H a_h^{\check{\theta}_h-1}$, the mean and the variance are given as follows:

$$\hat{\theta}_h = \langle \theta_h \rangle_{p(\boldsymbol{\theta}|\check{\boldsymbol{\theta}})} = \frac{\check{\theta}_h}{\check{\theta}_0}, \quad \langle (\theta_h - \hat{\theta}_h)^2 \rangle_{p(\boldsymbol{\theta}|\check{\boldsymbol{\theta}})} = \frac{\check{\theta}_h(\check{\theta}_0 - \check{\theta}_h)}{\check{\theta}_0^2(\check{\theta}_0 + 1)},$$

where $\check{\theta}_0 = \sum_{h=1}^H \check{\theta}_h$.

For fixed N, R , defined by Eq.(15), diverges to $+\infty$ if $\check{\Theta}_{m,h} \rightarrow +0$ for any (m, h) or $\check{B}_{l,h} \rightarrow +0$ for any (l, h) . Therefore, the global minimizer of the free energy (14) is in the interior of the domain, where the free energy is differentiable. Consequently, the global minimizer is a stationary point. The stationary condition (12) implies that

$$\check{\Theta}_{m,h} \geq \alpha, \quad \check{B}_{l,h} \geq \eta, \quad (23)$$

$$\sum_{h=1}^H \check{\Theta}_{m,h} = \sum_{h=1}^H \alpha + N^{(m)}, \quad \sum_{l=1}^L \check{B}_{l,h} = \sum_{l=1}^L \eta + \sum_{m=1}^M (\check{\Theta}_{m,h} - \alpha). \quad (24)$$

Therefore, we have

$$\langle (\Theta_{m,h} - \hat{\Theta}_{m,h})^2 \rangle_{q(\boldsymbol{\Theta})} = O_p(N^{-2}) \quad \text{for all } (m, h), \quad (25)$$

$$\left(\max_m \hat{\Theta}_{m,h} \right)^2 \langle (B_{l,h} - \hat{B}_{l,h})^2 \rangle_{q(\boldsymbol{B})} = O_p(N^{-2}) \quad \text{for all } (l, h), \quad (26)$$

which leads to Eq.(17).

By using Eq.(22), Q is bounded as follows:

$$\underline{Q} \leq Q \leq \overline{Q},$$

where

$$\begin{aligned} \overline{Q} &= - \sum_{m=1}^M N^{(m)} \sum_{l=1}^L V_{l,m} \log \left(\sum_{h=1}^H \frac{\check{\Theta}_{m,h}}{\sum_{h'=1}^H \check{\Theta}_{m,h'}} \frac{\check{B}_{l,h}}{\sum_{l'=1}^L \check{B}_{l',h}} \frac{\exp\left(-\frac{1}{\check{\Theta}_{m,h}}\right)}{\exp\left(-\frac{1}{2 \sum_{h'=1}^H \check{\Theta}_{m,h'}}\right)} \frac{\exp\left(-\frac{1}{\check{B}_{l,h}}\right)}{\exp\left(-\frac{1}{2 \sum_{l'=1}^L \check{B}_{l',h}}\right)} \right), \\ \underline{Q} &= - \sum_{m=1}^M N^{(m)} \sum_{l=1}^L V_{l,m} \log \left(\sum_{h=1}^H \frac{\check{\Theta}_{m,h}}{\sum_{h'=1}^H \check{\Theta}_{m,h'}} \frac{\check{B}_{l,h}}{\sum_{l'=1}^L \check{B}_{l',h}} \frac{\exp\left(-\frac{1}{2 \check{\Theta}_{m,h}}\right)}{\exp\left(-\frac{1}{\sum_{h'=1}^H \check{\Theta}_{m,h'}}\right)} \frac{\exp\left(-\frac{1}{2 \check{B}_{l,h}}\right)}{\exp\left(-\frac{1}{\sum_{l'=1}^L \check{B}_{l',h}}\right)} \right). \end{aligned}$$

Using Eqs.(25) and (26), we have Eq.(18), which completes the proof of Lemma 1. $\square$

B Proof of Lemma 3

By using the bounds (21) and (22), R can be bounded as

$$\underline{R} \leq R \leq \overline{R}, \quad (27)$$

where

$$\begin{aligned} \underline{R} &= - \sum_{m=1}^M \log \left(\frac{\Gamma(H\alpha)}{\Gamma(\alpha)^H} \right) - \sum_{h=1}^H \log \left(\frac{\Gamma(L\eta)}{\Gamma(\eta)} \right) - \frac{M(H-1) + H(L-1)}{2} \log(2\pi) \\ &\quad + \sum_{m=1}^M \left\{ \left(H\alpha - \frac{1}{2} \right) \log \sum_{h=1}^H \check{\Theta}_{m,h} - \sum_{h=1}^H \left(\alpha - \frac{1}{2} \right) \log \check{\Theta}_{m,h} \right\} \end{aligned}$$

$$\begin{aligned}
& + \sum_{h=1}^H \left\{ \left(L\eta - \frac{1}{2} \right) \log \sum_{l=1}^L \check{B}_{l,h} - \sum_{l=1}^L \left(\eta - \frac{1}{2} \right) \log \check{B}_{l,h} \right\} \\
& + \sum_{m=1}^M \left\{ - \sum_{h=1}^H \frac{1}{12\check{\Theta}_{m,h}} - \sum_{h=1}^H \left(\check{\Theta}_{m,h} - \alpha \right) \left(\frac{1}{\check{\Theta}_{m,h}} - \frac{1}{2 \sum_{h'=1}^H \check{\Theta}_{m,h'}} \right) \right\} \\
& + \sum_{h=1}^H \left\{ - \sum_{l=1}^L \frac{1}{12\check{B}_{l,h}} - \sum_{l=1}^L \left(\check{B}_{l,h} - \eta \right) \left(\frac{1}{\check{B}_{l,h}} - \frac{1}{2 \sum_{l'=1}^L \check{B}_{l',h}} \right) \right\}, \tag{28} \\
\bar{R} = & - \sum_{m=1}^M \log \left(\frac{\Gamma(H\alpha)}{\Gamma(\alpha)^H} \right) - \sum_{h=1}^H \log \left(\frac{\Gamma(L\eta)}{\Gamma(\eta)^L} \right) - \frac{M(H-1) + H(L-1)}{2} \log(2\pi) \\
& + \sum_{m=1}^M \left\{ \left(H\alpha - \frac{1}{2} \right) \log \sum_{h=1}^H \check{\Theta}_{m,h} - \sum_{h=1}^H \left(\alpha - \frac{1}{2} \right) \log \check{\Theta}_{m,h} \right\} \\
& + \sum_{h=1}^H \left\{ \left(L\eta - \frac{1}{2} \right) \log \sum_{l=1}^L \check{B}_{l,h} - \sum_{l=1}^L \left(\eta - \frac{1}{2} \right) \log \check{B}_{l,h} \right\} \\
& + \sum_{m=1}^M \left\{ \frac{1}{12 \sum_{h=1}^H \check{\Theta}_{m,h}} - \sum_{h=1}^H \left(\check{\Theta}_{m,h} - \alpha \right) \left(\frac{1}{2\check{\Theta}_{m,h}} - \frac{1}{\sum_{h'=1}^H \check{\Theta}_{m,h'}} \right) \right\} \\
& + \sum_{h=1}^H \left\{ \frac{1}{12 \sum_{l=1}^L \check{B}_{l,h}} - \sum_{l=1}^L \left(\check{B}_{l,h} - \eta \right) \left(\frac{1}{2\check{B}_{l,h}} - \frac{1}{\sum_{l'=1}^L \check{B}_{l',h}} \right) \right\}. \tag{29}
\end{aligned}$$

Eqs.(23) and (24) imply that

$$\begin{aligned}
R = & \sum_{m=1}^M \left\{ \left(H\alpha - \frac{1}{2} \right) \log \sum_{h=1}^H \check{\Theta}_{m,h} - \sum_{h=1}^H \left(\alpha - \frac{1}{2} \right) \log \check{\Theta}_{m,h} \right\} \\
& + \sum_{h=1}^H \left\{ \left(L\eta - \frac{1}{2} \right) \log \sum_{l=1}^L \check{B}_{l,h} - \sum_{l=1}^L \left(\eta - \frac{1}{2} \right) \log \check{B}_{l,h} \right\} + O_p(H(M+L)),
\end{aligned}$$

which leads to Eq.(20). This completes the proof of Lemma 3. $\square$

C Proof of Theorem 1

Lemma 2 and Lemma 3 imply that the free energy can be written as follows:

$$\begin{aligned}
F - S = & \left\{ M \left(H\alpha - \frac{1}{2} \right) + \hat{H} \left(L\eta - \frac{1}{2} \right) - \sum_{h=1}^{\hat{H}} \left(\widehat{M}^{(h)} \left(\alpha - \frac{1}{2} \right) + \widehat{L}^{(h)} \left(\eta - \frac{1}{2} \right) \right) \right\} \log N \\
& + (H - \hat{K}) \left(L\eta - \frac{1}{2} \right) \log L + O_p(\hat{J}N + LM). \tag{30}
\end{aligned}$$

Below, we investigate the leading term of the free energy (30) in different asymptotic limits.

In the limit when $N \rightarrow \infty$ with $L, M \sim O(1)$

In this case, the minimizer should satisfy

$$\widehat{\mathbf{B}} \widehat{\boldsymbol{\Theta}}^\top = \mathbf{B}^* \boldsymbol{\Theta}^{*\top} + O_p(N^{-1}), \tag{31}$$

making $\hat{J} = 0$ with probability 1, and the leading term of the order of $O_p(\log N)$:

$$F - S = \left\{ M \left(H\alpha - \frac{1}{2} \right) + \hat{H} \left(L\eta - \frac{1}{2} \right) - \sum_{h=1}^{\hat{H}} \left(\widehat{M}^{(h)} \left(\alpha - \frac{1}{2} \right) + \widehat{L}^{(h)} \left(\eta - \frac{1}{2} \right) \right) \right\} \log N + O_p(1).$$

Note that Eq.(31) implies the consistency of the predictive distribution.

In the limit when $N, M \rightarrow \infty$ with $\frac{M}{N}, L \sim O(1)$

In this case,

$$\hat{J} = o_p(\log N), \tag{32}$$

making the leading term of the order of $O_p(N \log N)$:

$$F - S = \left\{ M \left(H\alpha - \frac{1}{2} \right) - \sum_{h=1}^{\hat{H}} \widehat{M}^{(h)} \left(\alpha - \frac{1}{2} \right) \right\} \log N + o_p(N \log N).$$

Eq.(32) implies that the predictive distribution is not necessarily consistent— $o_p(\log N)$ components can deviate from the true matrix $\mathbf{B}^* \boldsymbol{\Theta}^{*\top}$ by the order of $O_p(1)$.

In the limit when $N, L \rightarrow \infty$ with $\frac{L}{N}, M \sim O(1)$

In this case, Eq.(32) holds, and the leading term of the free energy is of the order of $O_p(N \log N)$:

$$F - S = HL\eta \log N + o_p(N \log N).$$

In the limit when $N, L, M \rightarrow \infty$ with $\frac{L}{N}, \frac{M}{N} \sim O(1)$

In this case,

$$\hat{J} = o_p(N \log N), \quad (33)$$

and the leading term of the free energy is of the order of $O_p(N^2 \log N)$:

$$F - S = H(M\alpha + L\eta) \log N + o_p(N^2 \log N).$$

This completes the proof of Theorem 1. $\square$

D Proof of Corollary 1

From the compact representation when $\hat{H} = H^*, \hat{M}^{(h)} = M^{*(h)}$, and $\hat{L}^{(h)} = L^{*(h)}$, we can decompose a singular component into two, keeping $\hat{B}\hat{\Theta}^\top$ unchanged, so that

$$\hat{H} \rightarrow \hat{H} + 1, \quad (34)$$

$$\sum_{h=1}^H \hat{M}^{(h)} \rightarrow \sum_{h=1}^{H^*} \hat{M}^{(h)} + \Delta M \quad \text{for} \quad \min_h M^{*(h)} \leq \Delta M \leq \max_h M^{*(h)}, \quad (35)$$

$$\sum_{h=1}^H \hat{L}^{(h)} \rightarrow \sum_{h=1}^{H^*} \hat{L}^{(h)} + \Delta L \quad \text{for} \quad 0 \leq \Delta L \leq \max_h L^{*(h)}. \quad (36)$$

Here the lower-bound for ΔM in Eq.(35) corresponds to the case that the least frequent topic is chosen to be decomposed, while the upper-bound to the case that the most frequent topic is chosen. The lower-bound for ΔL in Eq.(36) corresponds to the decomposition such that some of the word-occurrences are moved to a new topic, while the upper-bound to the decomposition such that the topic with the widest vocabulary is copied to a new topic. Note that the bounds both for ΔM and ΔL are not always achievable simultaneously, when we choose one topic to decompose.

Below, we investigate the relation between the sparsity of the solution and the hyperparameter setting in different asymptotic limits.

In the limit when $N \rightarrow \infty$ with $L, M \sim O(1)$

The coefficient of the leading term of the free energy is

$$\lambda = M \left(H\alpha - \frac{1}{2} \right) + \sum_{h=1}^{\hat{H}} \left(L\eta - \frac{1}{2} - \hat{M}^{(h)} \left(\alpha - \frac{1}{2} \right) - \hat{L}^{(h)} \left(\eta - \frac{1}{2} \right) \right). \quad (37)$$

Note that the solution is sparse if Eq.(37) is increasing of $\hat{H}$, and dense if it is decreasing. Eqs.(34)–(36) imply the following:

1. When $0 < \eta \leq \frac{1}{2L}$ and $\alpha \leq \frac{1}{2}$, the solution is sparse if

$$L\eta - \frac{1}{2} - \min_h M^{*(h)} \left(\alpha - \frac{1}{2} \right) > 0, \text{ or equivalently, } \alpha < \frac{1}{2} - \frac{1}{\min_h M^{*(h)}} \left(\frac{1}{2} - L\eta \right),$$

and dense if

$$\alpha > \frac{1}{2} - \frac{1}{\min_h M^{*(h)}} \left(\frac{1}{2} - L\eta \right).$$

2. When $0 < \eta \leq \frac{1}{2L}$ and $\alpha > \frac{1}{2}$, the solution is sparse if

$$L\eta - \frac{1}{2} - \max_h M^{*(h)} \left(\alpha - \frac{1}{2} \right) > 0, \text{ or equivalently, } \alpha < \frac{1}{2} - \frac{1}{\max_h M^{*(h)}} \left(\frac{1}{2} - L\eta \right),$$

and dense if

$$\alpha > \frac{1}{2} - \frac{1}{\max_h M^{*(h)}} \left(\frac{1}{2} - L\eta \right).$$

Therefore, the solution is always dense in this case.

3. When $\frac{1}{2L} < \eta \leq \frac{1}{2}$ and $\alpha < \frac{1}{2}$, the solution is sparse if

$$L\eta - \frac{1}{2} - \min_h M^{*(h)} \left(\alpha - \frac{1}{2} \right) > 0, \text{ or equivalently, } \alpha < \frac{1}{2} + \frac{1}{\min_h M^{*(h)}} \left(L\eta - \frac{1}{2} \right),$$

and dense if

$$\alpha > \frac{1}{2} + \frac{1}{\min_h M^{*(h)}} \left(L\eta - \frac{1}{2} \right).$$

Therefore, the solution is always sparse in this case.

4. When $\frac{1}{2L} < \eta \leq \frac{1}{2}$ and $\alpha \geq \frac{1}{2}$, the solution is sparse if

$$L\eta - \frac{1}{2} - \max_h M^{*(h)} \left(\alpha - \frac{1}{2} \right) > 0, \text{ or equivalently, } \alpha < \frac{1}{2} + \frac{1}{\max_h M^{*(h)}} \left(L\eta - \frac{1}{2} \right),$$

and dense if

$$\alpha > \frac{1}{2} + \frac{1}{\max_h M^{*(h)}} \left(L\eta - \frac{1}{2} \right).$$

5. When $\eta > \frac{1}{2}$ and $\alpha < \frac{1}{2}$, the solution is sparse if

$$L\eta - \frac{1}{2} - \max_h \left(M^{*(h)} \left(\alpha - \frac{1}{2} \right) + L^{*(h)} \left(\eta - \frac{1}{2} \right) \right) > 0, \quad (38)$$

and dense if

$$L\eta - \frac{1}{2} - \max_h \left(M^{*(h)} \left(\alpha - \frac{1}{2} \right) + L^{*(h)} \left(\eta - \frac{1}{2} \right) \right) < 0. \quad (39)$$

Therefore, the solution is sparse if

$$L\eta - \frac{1}{2} - \min_h M^{*(h)} \left(\alpha - \frac{1}{2} \right) - \max_h L^{*(h)} \left(\eta - \frac{1}{2} \right) > 0,$$

$$\text{or equivalently, } \alpha < \frac{1}{2} + \frac{1}{\min_h M^{*(h)}} \left(L\eta - \frac{1}{2} - \max_h L^{*(h)} \left(\eta - \frac{1}{2} \right) \right),$$

and dense if

$$L\eta - \frac{1}{2} - \max_h M^{*(h)} \left(\alpha - \frac{1}{2} \right) - \max_h L^{*(h)} \left(\eta - \frac{1}{2} \right) < 0,$$

$$\text{or equivalently, } \alpha > \frac{1}{2} + \frac{1}{\max_h M^{*(h)}} \left(L\eta - \frac{1}{2} - \max_h L^{*(h)} \left(\eta - \frac{1}{2} \right) \right).$$

Therefore, the solution is always sparse in this case.

6. When $\eta > \frac{1}{2}$ and $\alpha \geq \frac{1}{2}$, the solution is sparse if Eq.(38) holds, and dense if Eq.(39) holds. Therefore, the solution is sparse if

$$L\eta - \frac{1}{2} - \max_h M^{*(h)} \left(\alpha - \frac{1}{2} \right) - \max_h L^{*(h)} \left(\eta - \frac{1}{2} \right) > 0,$$

$$\text{or equivalently, } \alpha < \frac{1}{2} + \frac{1}{\max_h M^{*(h)}} \left(L\eta - \frac{1}{2} - \max_h L^{*(h)} \left(\eta - \frac{1}{2} \right) \right),$$

and dense if

$$L\eta - \frac{1}{2} - \min_h M^{*(h)} \left(\alpha - \frac{1}{2} \right) - \max_h L^{*(h)} \left(\eta - \frac{1}{2} \right) < 0,$$

$$\text{or equivalently, } \alpha > \frac{1}{2} + \frac{1}{\min_h M^{*(h)}} \left(L\eta - \frac{1}{2} - \max_h L^{*(h)} \left(\eta - \frac{1}{2} \right) \right).$$

Thus, we can conclude that, in this case, the solution is sparse if

$$\alpha < \frac{1}{2} + \frac{L-1}{2 \max_h M^{*(h)}},$$

and dense if

$$\alpha > \frac{1}{2} + \frac{L\eta - \frac{1}{2}}{\min_h M^{*(h)}}.$$

Summarizing the above, we have the following lemma:

Lemma 4 When $0 < \eta \leq \frac{1}{2L}$, the solution is sparse if $\alpha < \frac{1}{2} - \frac{\frac{1}{2} - L\eta}{\min_h M^{*(h)}}$, and dense if $\alpha > \frac{1}{2} - \frac{\frac{1}{2} - L\eta}{\min_h M^{*(h)}}$. When $\frac{1}{2L} < \eta \leq \frac{1}{2}$, the solution is sparse if $\alpha < \frac{1}{2} + \frac{L\eta - \frac{1}{2}}{\max_h M^{*(h)}}$, and dense if $\alpha > \frac{1}{2} + \frac{L\eta - \frac{1}{2}}{\max_h M^{*(h)}}$. When $\eta > \frac{1}{2}$, the solution is sparse if $\alpha < \frac{1}{2} + \frac{L-1}{2 \max_h M^{*(h)}}$, and dense if $\alpha > \frac{1}{2} + \frac{L\eta - \frac{1}{2}}{\min_h M^{*(h)}}$.

In the limit when $N, M \rightarrow \infty$ with $\frac{M}{N}, L \sim O(1)$

The coefficient of the leading term of the free energy is given by

$$\lambda = M \left(H\alpha - \frac{1}{2} \right) - \sum_{h=1}^{\hat{H}} \widehat{M}^{(h)} \left(\alpha - \frac{1}{2} \right). \quad (40)$$

Although the predictive distribution does not necessarily converges to the true distribution, we can investigate the sparsity of the solution by considering the duplication rules (34)–(36) that keep $\widehat{\mathbf{B}}\widehat{\boldsymbol{\Theta}}^\top$ unchanged. It is clear that Eq.(40) is increasing of $\widehat{H}$ if $\alpha < \frac{1}{2}$, and decreasing if $\alpha > \frac{1}{2}$. Combing this result with Lemma 4 completes the proof of Corollary 1. $\square$

E Proof of Theorem 2

We analyze PBA learning, PBB learning, and MAP estimation, and then summarize the results.

E.1 PBA Learning

The free energy for PBA learning is given as follows:

$$F^{\text{PBA}} = \chi_B + R^{\text{PBA}} + Q^{\text{PBA}}, \quad (41)$$

where χ_B is a large constant corresponding to the negative entropy of the delta functions, and

$$\begin{aligned} R^{\text{PBA}} &= \left\langle \log \frac{q(\boldsymbol{\Theta})q(\mathbf{B})}{p(\boldsymbol{\Theta}|\alpha)p(\mathbf{B}|\eta)} \right\rangle_{q^{\text{PBA}}(\boldsymbol{\Theta}, \mathbf{B})} \\ &= \sum_{m=1}^M \left(\log \frac{\Gamma(\sum_{h=1}^H \check{\Theta}_{m,h}^{\text{PBA}})}{\prod_{h=1}^H \Gamma(\check{\Theta}_{m,h}^{\text{PBA}})} \frac{\Gamma(\alpha)^H}{\Gamma(H\alpha)} + \sum_{h=1}^H \left(\check{\Theta}_{m,h}^{\text{PBA}} - \alpha \right) \left(\Psi(\check{\Theta}_{m,h}^{\text{PBA}}) - \Psi(\sum_{h'=1}^H \check{\Theta}_{m,h'}^{\text{PBA}}) \right) \right) \\ &\quad + \sum_{h=1}^H \left(\log \frac{\Gamma(\eta)^L}{\Gamma(L\eta)} + \sum_{l=1}^L (1 - \eta) \left(\log(\check{B}_{l,h}^{\text{PBA}}) - \log(\sum_{l'=1}^L \check{B}_{l',h}^{\text{PBA}}) \right) \right), \end{aligned} \quad (42)$$

$$\begin{aligned} Q^{\text{PBA}} &= \left\langle \log \frac{q(\{\mathbf{z}^{(n,m)}\})}{p(\{\mathbf{w}^{(n,m)}\}, \{\mathbf{z}^{(n,m)}\}|\boldsymbol{\Theta}, \mathbf{B})} \right\rangle_{q^{\text{PBA}}(\boldsymbol{\Theta}, \mathbf{B}, \{\mathbf{z}^{(n,m)}\})} \\ &= - \sum_{m=1}^M N^{(m)} \sum_{l=1}^L V_{l,m} \log \left(\sum_{h=1}^H \frac{\exp(\Psi(\check{\Theta}_{m,h}^{\text{PBA}}))}{\exp(\Psi(\sum_{h'=1}^H \check{\Theta}_{m,h'}^{\text{PBA}}))} \frac{\check{B}_{l,h}^{\text{PBA}}}{\sum_{l'=1}^L \check{B}_{l',h}^{\text{PBA}}} \right). \end{aligned} \quad (43)$$

Let us first consider the case when $\eta < 1$. In this case, F diverges to $F \rightarrow -\infty$ with fixed N , when $\check{B}_{l,h} = O(1)$ for any (l, h) and $\check{B}_{l',h} \rightarrow +0$ for all other $l' \neq l$. Therefore, the solution is useless.

When $\eta \geq 1$, the solution satisfies the following stationary condition:

$$\check{\Theta}_{m,h}^{\text{PBA}} = \alpha + \sum_{n=1}^{N^{(m)}} \hat{z}_h^{\text{PBA}(n,m)}, \quad \check{B}_{l,h}^{\text{PBA}} = \eta - 1 + \sum_{m=1}^M \sum_{n=1}^{N^{(m)}} w_l^{(n,m)} \hat{z}_h^{\text{PBA}(n,m)}, \quad (44)$$

$$\hat{z}_h^{\text{PBA}(n,m)} = \frac{\exp(\Psi(\check{\Theta}_{m,h}^{\text{PBA}})) \prod_{l=1}^L (\check{B}_{l,h}^{\text{PBA}})^{w_l^{(n,m)}}}{\sum_{h'=1}^H \exp(\Psi(\check{\Theta}_{m,h'}^{\text{PBA}})) \prod_{l=1}^L (\check{B}_{l,h'}^{\text{PBA}})^{w_l^{(n,m)}}}. \quad (45)$$

In the same way as for VB learning, we can obtain the following lemma:

Lemma 5 Let $\widehat{\mathbf{B}}^{\text{PBA}} \widehat{\boldsymbol{\Theta}}^{\text{PBA}\top} = \langle \mathbf{B} \boldsymbol{\Theta}^\top \rangle_{q^{\text{PBA}}(\boldsymbol{\Theta}, \mathbf{B})}$. Then, it holds that

$$\langle (\mathbf{B} \boldsymbol{\Theta}^\top - \widehat{\mathbf{B}}^{\text{PBA}} \widehat{\boldsymbol{\Theta}}^{\text{PBA}\top})_{l,m}^2 \rangle_{q^{\text{PBA}}(\boldsymbol{\Theta}, \mathbf{B})} = O_p(N^{-2}), \quad (46)$$

$$Q^{\text{PBA}} = - \sum_{m=1}^M N^{(m)} \sum_{l=1}^L V_{l,m} \log(\widehat{\mathbf{B}}^{\text{PBA}} \widehat{\boldsymbol{\Theta}}^{\text{PBA}\top})_{l,m} + O_p(N^{-1}). \quad (47)$$

Q^{PBA} is minimized when $\widehat{\mathbf{B}}^{\text{PBA}} \widehat{\boldsymbol{\Theta}}^{\text{PBA}\top} = \mathbf{B}^* \boldsymbol{\Theta}^{*\top} + O_p(N^{-1})$, and it holds that

$$Q^{\text{PBA}} = S + O_p(\widehat{J}N + LM).$$

R^{PBA} is written as follows:

$$\begin{aligned} R^{\text{PBA}} &= \left\{ M \left(H\alpha - \frac{1}{2} \right) + \widehat{H}L(\eta - 1) - \sum_{h=1}^{\widehat{H}} \left(\widehat{M}^{(h)} \left(\alpha - \frac{1}{2} \right) + \widehat{L}^{(h)} (\eta - 1) \right) \right\} \log N \\ &\quad + (H - \widehat{H})L(\eta - 1) \log L + O_p(H(M + L)). \end{aligned} \quad (48)$$

Taking the different asymptotic limits, we obtain the following theorem:

Theorem 3 When $\eta < 1$, each column vector of $\hat{\mathbf{B}}^{\text{PBA}}$ has only one non-zero entry. Assume below that $\eta \geq 1$. In the limit when $N \rightarrow \infty$ with $L, M \sim O(1)$, it holds that $\hat{\mathcal{J}} = 0$ with probability 1, and

$$F^{\text{PBA}} = S + \left\{ M \left(H\alpha - \frac{1}{2} \right) + \hat{H}L(\eta - 1) - \sum_{h=1}^{\hat{H}} \left(\hat{M}^{(h)} \left(\alpha - \frac{1}{2} \right) + \hat{L}^{(h)} (\eta - 1) \right) \right\} \log N + O_p(1).$$

In the limit when $N, M \rightarrow \infty$ with $\frac{M}{N}, L \sim O(1)$, it holds that $\hat{\mathcal{J}} = o_p(\log N)$, and

$$F^{\text{PBA}} = S + \left\{ M \left(H\alpha - \frac{1}{2} \right) - \sum_{h=1}^{\hat{H}} \hat{M}^{(h)} \left(\alpha - \frac{1}{2} \right) \right\} \log N + o_p(N \log N).$$

In the limit when $N, L \rightarrow \infty$ with $\frac{L}{N}, M \sim O(1)$, it holds that $\hat{\mathcal{J}} = o_p(\log N)$, and

$$F^{\text{PBA}} = S + HL(\eta - 1) \log N + o_p(N \log N).$$

In the limit when $N, L, M \rightarrow \infty$ with $\frac{L}{N}, \frac{M}{N} \sim O(1)$, it holds that $\hat{\mathcal{J}} = o_p(N \log N)$, and

$$F^{\text{PBA}} = S + H(M\alpha + L(\eta - 1)) \log N + o_p(N^2 \log N).$$

Note that Theorem 3 provides no information on the sparsity of the PBA solution for $\eta < 1$. Below, we investigate the sparsity of the solution for $\eta \geq 1$.

In the limit when $N \rightarrow \infty$ with $L, M \sim O(1)$

The coefficient of the leading term of the free energy is

$$\lambda^{\text{PBA}} = M \left(H\alpha - \frac{1}{2} \right) + \sum_{h=1}^{\hat{H}} \left(L(\eta - 1) - \hat{M}^{(h)} \left(\alpha - \frac{1}{2} \right) - \hat{L}^{(h)} (\eta - 1) \right).$$

The solution is sparse if λ^{PBA} is increasing of $\hat{H}$, and dense if it is decreasing. We focus on the case when $\eta \geq 1$. Eqs.(34)–(36) imply the following:

1. When $\alpha < \frac{1}{2}$, the solution is sparse if

$$L(\eta - 1) - \max_h \left(M^{*(h)} \left(\alpha - \frac{1}{2} \right) + L^{*(h)} (\eta - 1) \right) > 0, \quad (49)$$

and dense if

$$L(\eta - 1) - \max_h \left(M^{*(h)} \left(\alpha - \frac{1}{2} \right) + L^{*(h)} (\eta - 1) \right) < 0. \quad (50)$$

Therefore, the solution is sparse if

$$L(\eta - 1) - \min_h M^{*(h)} \left(\alpha - \frac{1}{2} \right) - \max_h L^{*(h)} (\eta - 1) > 0,$$

or equivalently, $\alpha < \frac{1}{2} + \frac{(L - \max_h L^{*(h)}) (\eta - 1)}{\min_h M^{*(h)}}$,

and dense if

$$L(\eta - 1) - \max_h M^{*(h)} \left(\alpha - \frac{1}{2} \right) - \max_h L^{*(h)} (\eta - 1) < 0,$$

or equivalently, $\alpha > \frac{1}{2} + \frac{(L - \max_h L^{*(h)}) (\eta - 1)}{\max_h M^{*(h)}}$.

Therefore, the solution is always sparse in this case.

2. When $\alpha \geq \frac{1}{2}$, the solution is sparse if Eq.(49) holds, and dense if Eq.(50) holds. Therefore, the solution is sparse if

$$L(\eta - 1) - \max_h M^{*(h)} \left(\alpha - \frac{1}{2} \right) - \max_h L^{*(h)} (\eta - 1) > 0,$$

$$\text{or equivalently, } \alpha < \frac{1}{2} + \frac{(L - \max_h L^{*(h)}) (\eta - 1)}{\max_h M^{*(h)}},$$

and dense if

$$L(\eta - 1) - \min_h M^{*(h)} \left(\alpha - \frac{1}{2} \right) - \max_h L^{*(h)} (\eta - 1) < 0,$$

$$\text{or equivalently, } \alpha > \frac{1}{2} + \frac{(L - \max_h L^{*(h)}) (\eta - 1)}{\min_h M^{*(h)}}.$$

Thus, we can conclude that, in this case, the solution is sparse if

$$\alpha < \frac{1}{2},$$

and dense if

$$\alpha > \frac{1}{2} + \frac{L(\eta - 1)}{\min_h M^{*(h)}}.$$

Summarizing the above, we have the following lemma:

Lemma 6 Assume that $\eta \geq 1$. The solution is sparse if $\alpha < \frac{1}{2}$, and dense if $\alpha > \frac{1}{2} + \frac{L(\eta-1)}{\min_h M^{*(h)}}$.

In the limit when $N, M \rightarrow \infty$ with $\frac{M}{N}, L \sim O(1)$

The coefficient of the leading term of the free energy is given by

$$\lambda = M \left(H\alpha - \frac{1}{2} \right) - \sum_{h=1}^{\hat{H}} \widehat{M}^{(h)} \left(\alpha - \frac{1}{2} \right). \quad (51)$$

Although the predictive distribution does not necessarily converges to the true distribution, we can investigate the sparsity of the solution by considering the duplication rules (34)–(36) that keep $\widehat{\mathbf{B}}\widehat{\boldsymbol{\Theta}}^\top$ unchanged. It is clear that Eq.(51) is increasing of $\hat{H}$ if $\alpha < \frac{1}{2}$, and decreasing if $\alpha > \frac{1}{2}$. Combing this result with Lemma 6, we obtain the following corollary:

Corollary 2 Assume that $\eta \geq 1$. In the limit when $N \rightarrow \infty$ with $L, M \sim O(1)$, the PBA solution is sparse if $\alpha < \frac{1}{2}$, and dense if $\alpha > \frac{1}{2} + \frac{L(\eta-1)}{\min_h M^{*(h)}}$. In the limit when $N, M \rightarrow \infty$ with $\frac{M}{N}, L \sim O(1)$, the PBA solution is sparse if $\alpha < \frac{1}{2}$, and dense if $\alpha > \frac{1}{2}$.

E.2 PBB Learning

The free energy for PBB learning is given as follows:

$$F^{\text{PBB}} = \chi_\Theta + R^{\text{PBB}} + Q^{\text{PBB}}, \quad (52)$$

where χ_Θ is a large constant corresponding to the negative entropy of the delta functions, and

$$\begin{aligned} R^{\text{PBB}} &= \left\langle \log \frac{q(\boldsymbol{\Theta})q(\mathbf{B})}{p(\boldsymbol{\Theta}|\alpha)p(\mathbf{B}|\eta)} \right\rangle_{q^{\text{PBB}}(\boldsymbol{\Theta}, \mathbf{B})} \\ &= \sum_{m=1}^M \left(\log \frac{\Gamma(\alpha)^H}{\Gamma(H\alpha)} + \sum_{h=1}^H (1 - \alpha) \left(\log(\check{\Theta}_{m,h}^{\text{PBB}}) - \log(\sum_{h'=1}^H \check{\Theta}_{m,h'}^{\text{PBB}}) \right) \right) \\ &\quad + \sum_{h=1}^H \left(\log \frac{\Gamma(\sum_{l=1}^L \check{B}_{l,h}^{\text{PBB}})}{\prod_{l=1}^L \Gamma(\check{B}_{l,h}^{\text{PBB}})} \frac{\Gamma(\eta)^L}{\Gamma(L\eta)} + \sum_{l=1}^L \left(\check{B}_{l,h}^{\text{PBB}} - \eta \right) \left(\Psi(\check{B}_{l,h}^{\text{PBB}}) - \Psi(\sum_{l'=1}^L \check{B}_{l',h}^{\text{PBB}}) \right) \right), \end{aligned} \quad (53)$$

$$\begin{aligned}
Q^{\text{PBB}} &= \left\langle \log \frac{q(\{\mathbf{z}^{(n,m)}\})}{p(\{\mathbf{w}^{(n,m)}\}, \{\mathbf{z}^{(n,m)}\} | \boldsymbol{\Theta}, \mathbf{B})} \right\rangle_{q^{\text{PBB}}(\boldsymbol{\Theta}, \mathbf{B}, \{\mathbf{z}^{(n,m)}\})} \\
&= - \sum_{m=1}^M N^{(m)} \sum_{l=1}^L V_{l,m} \log \left(\sum_{h=1}^H \frac{\check{\Theta}_{m,h}^{\text{PBB}}}{\sum_{h'=1}^H \check{\Theta}_{m,h'}^{\text{PBB}}} \frac{\exp(\Psi(\check{B}_{l,h}^{\text{PBB}}))}{\exp(\Psi(\sum_{l'=1}^L \check{B}_{l',h}^{\text{PBB}}))} \right). \tag{54}
\end{aligned}$$

Let us first consider the case when $\alpha < 1$. In this case, F diverges to $F \rightarrow -\infty$ with fixed N , when $\check{\Theta}_{m,h} = O(1)$ for any (m, h) and $\check{\Theta}_{m,h'} \rightarrow +0$ for all other $h' \neq h$. Therefore, the solution is sparse (so sparse that the estimator is useless).

When $\alpha \geq 1$, the solution satisfies the following stationary condition:

$$\check{\Theta}_{m,h}^{\text{PBB}} = \alpha - 1 + \sum_{n=1}^{N^{(m)}} \hat{z}_h^{\text{PBB}(n,m)}, \quad \check{B}_{l,h}^{\text{PBB}} = \eta + \sum_{m=1}^M \sum_{n=1}^{N^{(m)}} w_l^{(n,m)} \hat{z}_h^{\text{PBB}(n,m)}, \tag{55}$$

$$\hat{z}_h^{\text{PBB}(n,m)} = \frac{\check{\Theta}_{m,h}^{\text{PBB}} \exp\left\{\sum_{l=1}^L w_l^{(n,m)} (\Psi(\check{B}_{l,h}^{\text{PBB}}) - \Psi(\sum_{l'=1}^L \check{B}_{l',h}^{\text{PBB}}))\right\}}{\sum_{h'=1}^H \check{\Theta}_{m,h'}^{\text{PBB}} \exp\left\{\sum_{l=1}^L w_l^{(n,m)} (\Psi(\check{B}_{l,h'}^{\text{PBB}}) - \Psi(\sum_{l'=1}^L \check{B}_{l',h'}^{\text{PBB}}))\right\}}. \tag{56}$$

In the same way as for VB and PBA learning, we can obtain the following lemma:

Lemma 7 Let $\hat{\mathbf{B}}^{\text{PBB}} \hat{\boldsymbol{\Theta}}^{\text{PBB}\top} = \langle \mathbf{B} \boldsymbol{\Theta}^\top \rangle_{q^{\text{PBB}}(\boldsymbol{\Theta}, \mathbf{B})}$. Then, it holds that

$$\langle (\mathbf{B} \boldsymbol{\Theta}^\top - \hat{\mathbf{B}}^{\text{PBB}} \hat{\boldsymbol{\Theta}}^{\text{PBB}\top})_{l,m}^2 \rangle_{q^{\text{PBB}}(\boldsymbol{\Theta}, \mathbf{B})} = O_p(N^{-2}), \tag{57}$$

$$Q^{\text{PBB}} = - \sum_{m=1}^M N^{(m)} \sum_{l=1}^L V_{l,m} \log(\hat{\mathbf{B}}^{\text{PBB}} \hat{\boldsymbol{\Theta}}^{\text{PBB}\top})_{l,m} + O_p(N^{-1}). \tag{58}$$

Q^{PBB} is minimized when $\hat{\mathbf{B}}^{\text{PBB}} \hat{\boldsymbol{\Theta}}^{\text{PBB}\top} = \mathbf{B}^* \boldsymbol{\Theta}^{*\top} + O_p(N^{-1})$, and it holds that

$$Q^{\text{PBB}} = S + O_p(\hat{J}N + LM).$$

R^{PBB} is written as follows:

$$\begin{aligned}
R^{\text{PBB}} &= \left\{ MH(\alpha - 1) + \hat{H} \left(L\eta - \frac{1}{2} \right) - \sum_{h=1}^{\hat{H}} \left(\widehat{M}^{(h)} (\alpha - 1) + \widehat{L}^{(h)} \left(\eta - \frac{1}{2} \right) \right) \right\} \log N \\
&\quad + (H - \hat{H}) \left(L\eta - \frac{1}{2} \right) \log L + O_p(H(M + L)). \tag{59}
\end{aligned}$$

Taking the different asymptotic limits, we obtain the following theorem:

Theorem 4 When $\alpha < 1$, each row vector of $\hat{\boldsymbol{\Theta}}^{\text{PBB}}$ has only one non-zero entry, and the PBB solution is sparse. Assume below that $\alpha \geq 1$. In the limit when $N \rightarrow \infty$ with $L, M \sim O(1)$, it holds that $\hat{J} = 0$ with probability 1, and

$$\begin{aligned}
F^{\text{PBB}} &= S + \left\{ MH(\alpha - 1) + \hat{H} \left(L\eta - \frac{1}{2} \right) - \sum_{h=1}^{\hat{H}} \left(\widehat{M}^{(h)} (\alpha - 1) + \widehat{L}^{(h)} \left(\eta - \frac{1}{2} \right) \right) \right\} \log N \\
&\quad + O_p(1).
\end{aligned}$$

In the limit when $N, M \rightarrow \infty$ with $\frac{M}{N}, L \sim O(1)$, it holds that $\hat{J} = o_p(\log N)$, and

$$F^{\text{PBB}} = S + \left\{ MH(\alpha - 1) - \sum_{h=1}^{\hat{H}} \widehat{M}^{(h)} (\alpha - 1) \right\} \log N + o_p(N \log N).$$

In the limit when $N, L \rightarrow \infty$ with $\frac{L}{N}, M \sim O(1)$, it holds that $\hat{J} = o_p(\log N)$, and

$$F^{\text{PBB}} = S + HL\eta \log N + o_p(N \log N).$$

In the limit when $N, L, M \rightarrow \infty$ with $\frac{L}{N}, \frac{M}{N} \sim O(1)$, it holds that $\hat{J} = o_p(N \log N)$, and

$$F^{\text{PBB}} = S + H(M(\alpha - 1) + L\eta) \log N + o_p(N^2 \log N).$$

Theorem 4 states that the PBB solution is sparse when $\alpha < 1$. Below, we investigate the sparsity of the solution for $\alpha \geq 1$.

In the limit when $N \rightarrow \infty$ with $L, M \sim O(1)$

The coefficient of the leading term of the free energy is

$$\lambda^{\text{PBB}} = MH(\alpha - 1) + \sum_{h=1}^{\widehat{H}} \left(L\eta - \frac{1}{2} - \widehat{M}^{(h)}(\alpha - 1) - \widehat{L}^{(h)}\left(\eta - \frac{1}{2}\right) \right).$$

The solution is sparse if λ^{PBB} is increasing of $\widehat{H}$, and dense if it is decreasing. We focus on the case when $\alpha \geq 1$. Eqs.(34)–(36) imply the following:

1. When $0 < \eta \leq \frac{1}{2L}$, the solution is sparse if

$$L\eta - \frac{1}{2} - \max_h M^{*(h)}(\alpha - 1) > 0, \text{ or equivalently, } \alpha < 1 - \frac{1}{\max_h M^{*(h)}} \left(\frac{1}{2} - L\eta \right),$$

and dense if

$$\alpha > 1 - \frac{1}{\max_h M^{*(h)}} \left(\frac{1}{2} - L\eta \right).$$

Therefore, the solution is always dense in this case.

2. When $\frac{1}{2L} < \eta \leq \frac{1}{2}$, the solution is sparse if

$$L\eta - \frac{1}{2} - \max_h M^{*(h)}(\alpha - 1) > 0, \text{ or equivalently, } \alpha < 1 + \frac{L\eta - \frac{1}{2}}{\max_h M^{*(h)}},$$

and dense if

$$\alpha > 1 + \frac{L\eta - \frac{1}{2}}{\max_h M^{*(h)}}.$$

3. When $\eta > \frac{1}{2}$, the solution is sparse if

$$L\eta - \frac{1}{2} - \max_h \left(M^{*(h)}(\alpha - 1) + L^{*(h)}\left(\eta - \frac{1}{2}\right) \right) > 0, \quad (60)$$

and dense if

$$L\eta - \frac{1}{2} - \max_h \left(M^{*(h)}(\alpha - 1) + L^{*(h)}\left(\eta - \frac{1}{2}\right) \right) < 0. \quad (61)$$

Therefore, the solution is sparse if

$$L\eta - \frac{1}{2} - \max_h M^{*(h)}(\alpha - 1) - \max_h L^{*(h)}\left(\eta - \frac{1}{2}\right) > 0,$$

$$\text{or equivalently, } \alpha < 1 + \frac{1}{\max_h M^{*(h)}} \left(L\eta - \frac{1}{2} - \max_h L^{*(h)}\left(\eta - \frac{1}{2}\right) \right),$$

and dense if

$$L\eta - \frac{1}{2} - \min_h M^{*(h)}(\alpha - 1) - \max_h L^{*(h)}\left(\eta - \frac{1}{2}\right) < 0,$$

$$\text{or equivalently, } \alpha > 1 + \frac{1}{\min_h M^{*(h)}} \left(L\eta - \frac{1}{2} - \max_h L^{*(h)}\left(\eta - \frac{1}{2}\right) \right).$$

Thus, we can conclude that, in this case, the solution is sparse if

$$\alpha < 1 + \frac{L - 1}{2 \max_h M^{*(h)}},$$

and dense if

$$\alpha > 1 + \frac{L\eta - \frac{1}{2}}{\min_h M^{*(h)}}.$$

Summarizing the above, we have the following lemma:

Lemma 8 Assume that $\alpha \geq 1$. When $0 < \eta \leq \frac{1}{2L}$, the solution is always dense. When $\frac{1}{2L} < \eta \leq \frac{1}{2}$, the solution is sparse if $\alpha < 1 + \frac{L\eta - \frac{1}{2}}{\max_h M^{*(h)}}$, and dense if $\alpha > 1 + \frac{L\eta - \frac{1}{2}}{\max_h M^{*(h)}}$. When $\eta > \frac{1}{2}$, the solution is sparse if $\alpha < 1 + \frac{L-1}{2 \max_h M^{*(h)}}$, and dense if $\alpha > 1 + \frac{L\eta - \frac{1}{2}}{\min_h M^{*(h)}}$.

In the limit when $N, M \rightarrow \infty$ with $\frac{M}{N}, L \sim O(1)$

The coefficient of the leading term of the free energy is given by

$$\lambda = M(H\alpha - 1) - \sum_{h=1}^{\hat{H}} \widehat{M}^{(h)}(\alpha - 1). \quad (62)$$

Although the predictive distribution does not necessarily converges to the true distribution, we can investigate the sparsity of the solution by considering the duplication rules (34)–(36) that keep $\widehat{\mathbf{B}}\widehat{\boldsymbol{\Theta}}^\top$ unchanged. It is clear that Eq.(62) is decreasing of $\hat{H}$ if $\alpha > 1$. Combing this result with Theorem 4, which states that the PBB solution is sparse when $\alpha < 1$, and Lemma 8, we obtain the following corollary:

Corollary 3 Consider the limit when $N \rightarrow \infty$ with $L, M \sim O(1)$. When $0 < \eta \leq \frac{1}{2L}$, the PBB solution is sparse if $\alpha < 1$, and dense if $\alpha > 1$. When $\frac{1}{2L} < \eta \leq \frac{1}{2}$, the PBB solution is sparse if $\alpha < 1 + \frac{L\eta - \frac{1}{2}}{\max_h M^{*(h)}}$, and dense if $\alpha > 1 + \frac{L\eta - \frac{1}{2}}{\max_h M^{*(h)}}$. When $\eta > \frac{1}{2}$, the PBB solution is sparse if $\alpha < 1 + \frac{L-1}{2\max_h M^{*(h)}}$, and dense if $\alpha > 1 + \frac{L\eta - \frac{1}{2}}{\min_h M^{*(h)}}$. In the limit when $N, M \rightarrow \infty$ with $\frac{M}{N}, L \sim O(1)$, the PBB solution is sparse if $\alpha < 1$, and dense if $\alpha > 1$.

E.3 MAP Learning

The free energy for MAP learning is given as follows:

$$F^{\text{MAP}} = \chi_{\boldsymbol{\Theta}} + \chi_{\mathbf{B}} + R^{\text{MAP}} + Q^{\text{MAP}}, \quad (63)$$

where $\chi_{\boldsymbol{\Theta}}$ and $\chi_{\mathbf{B}}$ are large constants corresponding to the negative entropies of the delta functions, and

$$\begin{aligned} R^{\text{MAP}} &= \left\langle \log \frac{q(\boldsymbol{\Theta})q(\mathbf{B})}{p(\boldsymbol{\Theta}|\alpha)p(\mathbf{B}|\eta)} \right\rangle_{q^{\text{MAP}}(\boldsymbol{\Theta}, \mathbf{B})} \\ &= \sum_{m=1}^M \left(\log \frac{\Gamma(\alpha)^H}{\Gamma(H\alpha)} + \sum_{h=1}^H (1 - \alpha) \left(\log(\check{\Theta}_{m,h}^{\text{MAP}}) - \log(\sum_{h'=1}^H \check{\Theta}_{m,h'}^{\text{MAP}}) \right) \right) \\ &\quad + \sum_{h=1}^H \left(\log \frac{\Gamma(\eta)^L}{\Gamma(L\eta)} + \sum_{l=1}^L (1 - \eta) \left(\log(\check{B}_{l,h}^{\text{MAP}}) - \log(\sum_{l'=1}^L \check{B}_{l',h}^{\text{MAP}}) \right) \right), \end{aligned} \quad (64)$$

$$\begin{aligned} Q^{\text{MAP}} &= \left\langle \log \frac{q(\{\mathbf{z}^{(n,m)}\})}{p(\{\mathbf{w}^{(n,m)}\}, \{\mathbf{z}^{(n,m)}\}|\boldsymbol{\Theta}, \mathbf{B})} \right\rangle_{q^{\text{MAP}}(\boldsymbol{\Theta}, \mathbf{B}, \{\mathbf{z}^{(n,m)}\})} \\ &= - \sum_{m=1}^M N^{(m)} \sum_{l=1}^L V_{l,m} \log \left(\sum_{h=1}^H \frac{\check{\Theta}_{m,h}^{\text{MAP}}}{\sum_{h'=1}^H \check{\Theta}_{m,h'}^{\text{MAP}}} \frac{\check{B}_{l,h}^{\text{MAP}}}{\sum_{l'=1}^L \check{B}_{l',h}^{\text{MAP}}} \right). \end{aligned} \quad (65)$$

Let us first consider the case when $\alpha < 1$. In this case, F diverges to $F \rightarrow -\infty$ with fixed N , when $\check{\Theta}_{m,h} = O(1)$ for any (h, m) and $\check{\Theta}_{m,h'} \rightarrow +0$ for all other $h' \neq h$. Therefore, the solution is sparse (so sparse that the estimator is useless). Similarly, assume that $\eta < 1$. Then, F diverges to $F \rightarrow -\infty$ with fixed N , when $\check{B}_{l,h} = O(1)$ for any (l, h) and $\check{B}_{l',h} \rightarrow +0$ for all other $l' \neq l$. Therefore, the solution is useless.

When $\alpha \geq 1$ and $\eta \geq 1$, the solution satisfies the following stationary condition:

$$\check{\Theta}_{m,h}^{\text{MAP}} = \alpha - 1 + \sum_{n=1}^{N^{(m)}} \widehat{z}_h^{\text{MAP}(n,m)}, \quad \check{B}_{l,h}^{\text{MAP}} = \eta - 1 + \sum_{m=1}^M \sum_{n=1}^{N^{(m)}} w_l^{(n,m)} \widehat{z}_h^{\text{MAP}(n,m)}, \quad (66)$$

$$\widehat{z}_h^{\text{MAP}(n,m)} = \frac{\check{\Theta}_{m,h}^{\text{MAP}} \prod_{l=1}^L (\check{B}_{l,h}^{\text{MAP}})^{w_l^{(n,m)}}}{\sum_{h'=1}^H \left(\check{\Theta}_{m,h'}^{\text{MAP}} \prod_{l=1}^L (\check{B}_{l,h'}^{\text{MAP}})^{w_l^{(n,m)}} \right)}. \quad (67)$$

In the same way as for VB, PBA, and PBB learning, we can obtain the following lemma:

Lemma 9 Let $\hat{\mathbf{B}}^{\text{MAP}} \hat{\boldsymbol{\Theta}}^{\text{MAP}\top} = \langle \mathbf{B} \boldsymbol{\Theta}^\top \rangle_{q^{\text{MAP}}(\boldsymbol{\Theta}, \mathbf{B})}$. Then, Q^{MAP} is minimized when $\hat{\mathbf{B}}^{\text{MAP}} \hat{\boldsymbol{\Theta}}^{\text{MAP}\top} = \mathbf{B}^* \boldsymbol{\Theta}^{*\top} + O_p(N^{-1})$, and it holds that

$$Q^{\text{MAP}} = S + O_p(\hat{J}N + LM).$$

R^{MAP} is written as follows:

$$\begin{aligned} R^{\text{MAP}} = & \left\{ MH(\alpha - 1) + \hat{H}L(\eta - 1) - \sum_{h=1}^{\hat{H}} \left(\widehat{M}^{(h)}(\alpha - 1) + \widehat{L}^{(h)}(\eta - 1) \right) \right\} \log N \\ & + (H - \hat{H})L(\eta - 1) \log L + O_p(H(M + L)). \end{aligned} \quad (68)$$

Taking the different asymptotic limits, we obtain the following theorem:

Theorem 5 When $\alpha < 1$, each row vector of $\hat{\boldsymbol{\Theta}}^{\text{MAP}}$ has only one non-zero entry, and the MAP solution is sparse. When $\eta < 1$, each column vector of $\hat{\mathbf{B}}^{\text{MAP}}$ has only one non-zero entry. Assume below that $\alpha, \eta \geq 1$. In the limit when $N \rightarrow \infty$ with $L, M \sim O(1)$, it holds that $\hat{J} = 0$ with probability 1, and

$$\begin{aligned} F^{\text{MAP}} = & S + \left\{ MH(\alpha - 1) + \hat{H}L(\eta - 1) - \sum_{h=1}^{\hat{H}} \left(\widehat{M}^{(h)}(\alpha - 1) + \widehat{L}^{(h)}(\eta - 1) \right) \right\} \log N \\ & + O_p(1). \end{aligned}$$

In the limit when $N, M \rightarrow \infty$ with $\frac{M}{N}, L \sim O(1)$, it holds that $\hat{J} = o_p(\log N)$, and

$$F^{\text{MAP}} = S + \left\{ MH(\alpha - 1) - \sum_{h=1}^{\hat{H}} \widehat{M}^{(h)}(\alpha - 1) \right\} \log N + o_p(N \log N).$$

In the limit when $N, L \rightarrow \infty$ with $\frac{L}{N}, M \sim O(1)$, it holds that $\hat{J} = o_p(\log N)$, and

$$F^{\text{MAP}} = S + HL(\eta - 1) \log N + o_p(N \log N).$$

In the limit when $N, L, M \rightarrow \infty$ with $\frac{L}{N}, \frac{M}{N} \sim O(1)$, it holds that $\hat{J} = o_p(N \log N)$, and

$$F^{\text{MAP}} = S + H(M(\alpha - 1) + L(\eta - 1)) \log N + o_p(N^2 \log N).$$

Theorem 5 states that the MAP solution is sparse when $\alpha < 1$. However it provides no information on the sparsity of the MAP solution for $\eta < 1$. Below, we investigate the sparsity of the solution for $\alpha, \eta \geq 1$.

In the limit when $N \rightarrow \infty$ with $L, M \sim O(1)$

The coefficient of the leading term of the free energy is

$$\lambda^{\text{MAP}} = MH(\alpha - 1) + \sum_{h=1}^{\hat{H}} \left(L(\eta - 1) - \widehat{M}^{(h)}(\alpha - 1) - \widehat{L}^{(h)}(\eta - 1) \right).$$

The solution is sparse if λ^{MAP} is increasing of $\hat{H}$, and dense if it is decreasing. We focus on the case when $\alpha, \eta \geq 1$. Eqs.(34)–(36) imply the following:

The solution is sparse if

$$L(\eta - 1) - \max_h \left(M^{*(h)}(\alpha - 1) + L^{*(h)}(\eta - 1) \right) > 0, \quad (69)$$

and dense if

$$L(\eta - 1) - \max_h \left(M^{*(h)}(\alpha - 1) + L^{*(h)}(\eta - 1) \right) < 0. \quad (70)$$

Therefore, the solution is sparse if

$$L(\eta - 1) - \max_h M^{*(h)}(\alpha - 1) - \max_h L^{*(h)}(\eta - 1) > 0,$$

$$\text{or equivalently, } \alpha < 1 + \frac{(L - \max_h L^{*(h)})(\eta - 1)}{\max_h M^{*(h)}},$$

ad dense if

$$L(\eta - 1) - \min_h M^{*(h)} (\alpha - 1) - \max_h L^{*(h)} (\eta - 1) < 0,$$

or equivalently, $\alpha > 1 + \frac{(L - \max_h L^{*(h)})(\eta - 1)}{\min_h M^{*(h)}}.$

Thus, we can conclude that the solution is sparse if

$$\alpha < 1,$$

and dense if

$$\alpha > 1 + \frac{L(\eta - 1)}{\min_h M^{*(h)}}.$$

Summarizing the above, we have the following lemma:

Lemma 10 Assume that $\eta \geq 1$. The solution is sparse if $\alpha < 1$, and dense if $\alpha > 1 + \frac{L(\eta-1)}{\min_h M^{*(h)}}.$

In the limit when $N, M \rightarrow \infty$ with $\frac{M}{N}, L \sim O(1)$

The coefficient of the leading term of the free energy is given by

$$\lambda^{\text{MAP}} = MH (\alpha - 1) - \sum_{h=1}^{\hat{H}} \widehat{M}^{(h)} (\alpha - 1). \quad (71)$$

Although the predictive distribution does not necessarily converges to the true distribution, we can investigate the sparsity of the solution by considering the duplication rules (34)–(36) that keep $\widehat{\mathbf{B}}\widehat{\boldsymbol{\Theta}}^\top$ unchanged. It is clear that Eq.(71) is decreasing of $\widehat{H}$ if $\alpha > 1$. Combing this result with Theorem 5, which states that the MAP solution is sparse if $\alpha < 1$, and Lemma 10, we obtain the following corollary:

Corollary 4 Assume that $\eta \geq 1$. In the limit when $N \rightarrow \infty$ with $L, M \sim O(1)$, the MAP solution is sparse if $\alpha < 1$, and dense if $\alpha > 1 + \frac{L(\eta-1)}{\min_h M^{*(h)}}.$ In the limit when $N, M \rightarrow \infty$ with $\frac{M}{N}, L \sim O(1)$, the MAP solution is sparse if $\alpha < 1$, and dense if $\alpha > 1$.

E.4 Summary of Results

Summarizing Corollary 1, Corollary 2, Corollary 3, and Corollary 4 completes the proof of Theorem 2. $\square$