

# Bayesian Inference and Prediction of a Multiple-Change-Point Panel Model with Nonparametric Priors

Mark Fisher and Mark J. Jensen

Working Paper 2018-2a  
June 2018

**Abstract:** Change point models using hierarchical priors share in the information of each regime when estimating the parameter values of a regime. Because of this sharing, hierarchical priors have been very successful when estimating the parameter values of short-lived regimes and predicting the out-of-sample behavior of the regime parameters. However, the hierarchical priors have been parametric. Their parametric nature leads to global shrinkage that biases the estimates of the parameter coefficient of extraordinary regimes toward the value of the average regime. To overcome this shrinkage, we model the hierarchical prior nonparametrically by letting the hyperparameter's prior—in other words, the hyperprior—be unknown and modeling it with a Dirichlet processes prior. To apply a nonparametric hierarchical prior to the probability of a break occurring, we extend the change point model to a multiple-change-point panel model. The hierarchical prior then shares in the cross-sectional information of the break processes to estimate the transition probabilities. We apply our multiple-change-point panel model to a longitudinal data set of actively managed, U.S. equity, mutual fund returns to measure fund performance and investigate the chances of a skilled fund being skilled in the future.

JEL classification: C11, C14, C41, G11, G17

Key words: Bayesian nonparametric analysis, change points, Dirichlet process, hierarchical priors, mutual fund performance

<https://doi.org/10.29338/wp2018-02>

---

The authors thank Vikas Agarwal, Sylvia Frühwirth-Schnatter, and Paula Tkac for helping to make this paper better. They also thank the conference participants at the 2015 Bayesian Workshop of the Rimini Centre of Economic Analysis in Rimini, Italy; the 2015 European Seminar on Bayesian Econometrics in Gerzensee, Switzerland; the 2015 Conference on Computational and Financial Econometrics in London; the 2016 All Georgia Finance Conference in Atlanta, Georgia; the 2017 XVII ESTE Brazilian School of Time Series and Econometrics in São Carlos, Brazil; the department members of the Vienna University of Economics and Business, where one of the authors visited; McMaster University; the University of Montreal; Clemson University; and the University of North Carolina-Charlotte for their helpful comments and suggestions. The views expressed here are the authors' and not necessarily those of the Federal Reserve Bank of Atlanta or the Federal Reserve System. Any remaining errors are the authors' responsibility.

Please address questions regarding content to Mark Fisher, Research Department, Federal Reserve Bank of Atlanta, 1000 Peachtree Street NE, Atlanta, GA 30309-4470, 404-498-8757, [mark@markfisher.net](mailto:mark@markfisher.net), or Mark J. Jensen, Research Department, Federal Reserve Bank of Atlanta, 1000 Peachtree Street NE, Atlanta, GA 30309-4470, 404-498-8019, [mark.jensen@atl.frb.org](mailto:mark.jensen@atl.frb.org).

Federal Reserve Bank of Atlanta working papers, including revised versions, are available on the Atlanta Fed's website at [www.frbatlanta.org](http://www.frbatlanta.org). Click "Publications" and then "Working Papers." To receive e-mail notifications about new papers, use [frbatlanta.org/forms/subscribe](http://frbatlanta.org/forms/subscribe).

# 1 Introduction

In economics and finance there are many instances where longitudinal data changes regimes at different points in time when experiencing structural breaks. For example, models of the real gross domestic products of the world’s economies, or the returns from the group of actively managed mutual funds, can have parameters that change over time and across the cross-section. When measuring the performance of actively managed mutual funds it is the fund-regime performance parameter, alpha, that changes from regime to regime and from fund to fund. Allowing for changes in each fund’s alpha at different points in time allows there to be periods when some funds are skilled and other times when the same funds are unskilled. To model this regime parameter behavior over a cross-section of individuals our first contribution is to extend the change point model to a panel of change point processes.<sup>1</sup> In particular, we model the performance of actively managed mutual funds with a multiple-change-point panel model where the skill level of each fund follows its own change-point process.

Hierarchical priors have been an important advancement to the estimation of the regime parameters of structural break models (see Pesaran et al. (2006), Koop & Potter (2007), Geweke & Jiang (2011), Maheu & Song (2014) and Song (2014)). Under hierarchical priors the hyperparameters are unknown and their values are learned about through the information in the regimes. By mixing the conditional hierarchical prior over the posterior of its hyperparameters, the hierarchical prior’s posterior predictive distribution flexibly models how the parameters are distributed over the regimes.

These hierarchical priors, however, are restrictive in the sense that the prior for the hyperparameter, the hyperprior distribution, is assumed to be known. Assuming the hyperprior distribution is known can lead to the undesirable outcome of the regime parameter estimates shrinking towards the global average of the regimes (see Gelman et al. (2013), Chapter 5). For example, when measuring mutual fund skill the performance estimate for a highly skilled mutual fund shrinks to that of an ordinary fund (see Jones & Shanken (2005)).

To provide a better and more robust estimate of the parameters we model the hierarchical prior nonparametrically by letting the hyperprior be unknown, giving it a Dirichlet process prior, and estimating it along with the hyperparameters. Nonparametric hierarchical priors span the space of distributions, including multimodal, skewed, and kurtotic

---

<sup>1</sup>Frühwirth-Schnatter & Kaufmann (2006) model a panel of bank lending series with a multiple-structural-break model. Billio et al. (2016) also propose a panel of Markov switching vector autoregressive models but their panel is relatively small compared to the size of the panel we envision here.

distributions (see Lijoi et al. (2005)). This spanning property is especially important when estimating the parameters of an extraordinary individual or for an extreme regime. Our second contribution relaxes the parametric nature of the hierarchical priors and flexibly models the hierarchical priors nonparametrically by letting the hyperprior be unknown.

The last contribution of the paper is found by applying our multiple-change-point panel model to a cross-section of mutual fund return data and estimating the alphas for each fund and their different regimes. Estimates of these fund-regime alphas help answer if above market returns by a fund today are any indication of future excellence by the fund.<sup>2</sup> The answer depends on a number of factors, for instance, who the fund manager is, whether the fund changes its strategy or objective, if the fund adjusts its risk exposure, what the fund's flow of assets under management is, and how the fund manager is compensated. A change in any one of these factors can result in a new regime and a different value for alpha. We find there to be overwhelming empirical evidence in favor of mutual fund performance following a change-point process so there are no guarantees future performance will mirror today's.

Our plan for the paper is to describe in Section 2 our multiple-change-point panel model. Section 3 then defines the nonparametric hierarchical priors for the model parameters. In that section we also define the Dirichlet process prior for an unknown hyperprior distribution. A Markov chain Monte Carlo (MCMC) sampler of the unknowns is then described in Section 4. Section 4.4 explains how the posterior draws from the MCMC sampler are used to infer how the unknowns are distributed over mutual funds and their regimes. In Section 5 we apply the multiple-change-point panel model, using both nonparametric and parametric, hierarchical, priors, to a unbalanced panel of actively managed, US equity, mutual fund return data. Section 6 concludes with a summary of our findings.

## 2 Modeling a panel of multiple-change-point processes

Our multiple-change-point panel model extends the models of Chib (1998), Pesaran et al. (2006), Koop & Potter (2007) and Geweke & Jiang (2011). Instead of there being a single change point process, our model consists of a cross-section of multiple-change-point processes. Given our interest in measuring the empirical performance of mutual funds and their ability to sustain above market returns we describe the model in terms of mutual fund returns.

---

<sup>2</sup>See Grinblatt & Sheridan (1992), Hendricks et al. (1993), Goetzmann & Ibbotson (1994), Elton et al. (1996), Bollen & Busse (2004), and Busse & Irvine (2006) for empirical evidence favoring a mutual fund performance today predicting future performance and Carhart (1997) for a contrary view.

Let  $y_{i,t}$ ,  $i = 1, \dots, J$ , and,  $t = \tau_i, \dots, T_i$ , where  $1 \leq \tau_i$ , be the risk-free adjusted, gross returns in month  $t$  for the  $i$ th mutual fund from a longitudinal data set where the length of the series is  $\mathcal{T}_i = T_i - \tau_i + 1$ . In contrast to the existing literature on change-point models, the lengths of the time series are all small relative to the size of the cross-section,  $J$ . In addition, the panel does not have to be balanced. As a result our model can handle return histories,  $Y_{i,T_i} = (y_{i,\tau_i}, \dots, y_{i,T_i})'$ , that are of different lengths; ie.,  $\mathcal{T}_i$  does not have to equal  $\mathcal{T}_{i'}$ . The model can also handle return histories that do not line up in time so  $\tau_i$  does not need to be equal to  $\tau_{i'}$ .

Following Chib (1998) we model the change point processes with a hidden Markov chain process. Let  $M_i$  be the maximum number of regimes the  $i$ th mutual fund can experience. In addition, let  $m \in \{1, \dots, M_i\}$  index the fund-regimes where the probability of switching to a new fund-regime is  $q_i$ . Fund  $i$ 's restricted Markov break process at observation  $t$  is denoted by  $s_{i,t}$ , where  $t = \tau_i, \tau_i + 1, \dots$ . The change point process can then be represented by the following hierarchical prior

$$Pr(s_{i,t+1} = l | s_{i,t} = m, q_i) = \begin{cases} 1 - q_i & l = m < M_i, \\ q_i & l = m + 1, \\ 1 & l = m = M_i, \\ 0 & \text{otherwise,} \end{cases} \quad (1)$$

$Pr(s_{i,\tau_i} = 1) = 1$ , and where the next level in the hierarchy is the prior for  $q_i$ .

According to Eq. (1) once a break occurs a fund cannot return to an earlier regime, nor can it skip over a regime. Such change point behavior is not restrictive when measuring the persistence in mutual fund performance. Mutual fund skill and stock picking strategies are difficult to replicate over different time periods. Hence, even if a past trading strategy could be adopted, or a known skilled manager were to be hired, performance today would most likely differ from the past given the additional information available today and the changing market environment.

In Eq. (1), once  $s_{i,t}$  reaches the  $M_i$ th regime it stays there and cannot experience another break. To avoid a pile-up on  $M_i$  and the problems found in Koop & Potter (2007), we follow Bauwens et al. (2015) and set  $M_i$  apriori and monitor  $s_{i,t}$  to ensure that it never reaches  $M_i$  before  $t$  equals  $T_i$ . In addition, to speed up estimation we assign smaller  $M_i$ s to those funds with shorter histories,  $\mathcal{T}_i$  (see Fig. 1 for a graph of the value for the  $M_i$ s). This reduces the need to estimate a large number of out-of-sample, fund-regime parameters.<sup>3</sup>

---

<sup>3</sup>Our approach requires estimating the parameter values for all  $M_i$  regimes and evaluating their likelihoods. It would be computationally less demanding if we could use the approach of Maheu & Song (2014) and let the number of change points be unknown. However, this would require the priors for the regime parameters to be joint conjugate which would restrict the flexibility of our nonparametric hierarchical priors.

The transition probability,  $q_i$ , is the same as in Chib (1998) and Pesaran et al. (2006), in that it is constant over time. Notice also that  $q_i$  does not depend on the regime  $m$  (see McCulloch & Tsay (1993) and Geweke & Jiang (2011) for models with the same assumption). Each fund's transition probability is different but we assume they all come from the same cross-sectional distribution,  $q_i \stackrel{iid}{\sim} \pi_q$ ,  $i = 1, \dots, J$ .

As pointed out by Koop & Potter (2007), letting  $q_i$  be constant has consequences. It means the duration of a regime follows a Geometric distribution where shorter length regimes have a higher probability of occurrence than do longer regimes. Although restrictive, empirically we find that by implicitly modeling the duration distribution as an unknown distribution through an unknown  $\pi_q$ , the probability of a long-lived regime occurring is not that much lower than the probability of a short-lived regime.

To simplify our exposition on estimating our multiple-change-point panel model, we limit the regime parameters to the skill level and the variance of the mutual fund returns in a zero risk-factor model. Later, in our empirical analysis, we let the returns follow a four factor risk model where the unknown beta coefficients follow the same change point process underlying the regime parameter values of skill and variance.

Let  $\alpha_{im}$  and  $\sigma_{im}^2$  be the level of skill and variance of the  $m$ th regime for the  $i$ th fund. Using a dynamic structural investment model of mutual fund performance, Koijen (2014) shows that changes in a fund's strategy or its manager is equivalent to  $\alpha_{im}$  changing value with each regime  $m$ . When  $s_{i,t} = m$ , the conditional sampling distribution of  $y_{i,t}$  is then defined to be

$$p(y_{i,t}|Y_{i,t-1}, s_{i,t} = m) \equiv N(\alpha_{im}, \sigma_{im}^2). \quad (2)$$

We could replace normally distributed returns with the scaled normal mixture representation of a Student-t distribution. A time-varying stochastic volatility process for the variances could also be used in place of the constant fund-regime variance. However, monthly mutual fund returns do not exhibit the leptokurtosis and time-varying volatility prevalent in daily stock return data. Instead, monthly mutual fund returns are known for their homoskedasticity and normality.

Let  $\mathcal{Y} := \{Y_{i,T_i}\}_{i=1}^J$  be the entire panel of mutual fund return data, and  $\boldsymbol{\alpha}_i = (\alpha_{i1}, \dots, \alpha_{iM_i})'$ , and,  $\boldsymbol{\sigma}_i^2 = (\sigma_{i1}^2, \dots, \sigma_{iM_i}^2)'$ , be all the regime parameters for the  $i$ th mutual fund. Denote the history of the latent change points up to time  $t$  as  $S_{it} = (s_{i,\tau_i}, \dots, s_{i,t})$  so that  $S_{iT_i}$  contains the entire in-sample history of the change points. If we assume trading strategies and stock picking skills are proprietary secrets, and, therefore, cannot be copied by, or transferred to,

another fund, the joint likelihood is written as

$$p(\mathcal{Y}|\boldsymbol{\alpha}, \boldsymbol{\sigma}^2, \mathcal{S}) = \prod_{i=1}^J \prod_{t=\tau_i}^{T_i} p(y_{i,t}|\alpha_{is_{i,t}}, \sigma_{is_{i,t}}^2), \quad (3)$$

where  $\boldsymbol{\alpha} := \{\boldsymbol{\alpha}_i\}_{i=1}^J$ ,  $\boldsymbol{\sigma}^2 := \{\boldsymbol{\sigma}_i^2\}_{i=1}^J$ , and  $\mathcal{S} := \{S_{iT_i}\}_{i=1}^J$ . This factoring of the likelihood simplifies both the exposition and the computation, but it is not necessary. One could choose to follow the approach of Jones & Shanken (2005) and model the cross-correlation between fund returns with a latent residual factor structure.

### 3 Nonparametric hierarchical priors

What one assumes about the priors for the regime parameters and transition probabilities plays an important role in the behavior of the multiple-change-point panel model and the measurement of mutual fund performance and its persistence. Specification of the prior can affect the expectation about the future performance of the actively managed mutual fund industry, the level of certainty in a fund's measure of performance, and the inference about how skill, risk taking, and the probability of a new regime, are distributed across the population of mutual funds and their regimes.

To answer these and other mutual fund performance related questions with our multiple-change-point panel model we denote the prior distribution for  $\alpha_{im}$ ,  $\sigma_{im}^2$ , and,  $q_i$ , as  $\pi_\alpha$ ,  $\pi_\sigma$ , and  $\pi_q$ , respectively. We also assume each prior is unknown and modeled with a nonparametric hierarchical prior where the hyperprior is unknown and follows the Dirichlet process (DP) of Ferguson (1973).

Our approach is similar but slightly different from Frühwirth-Schnatter & Kaufmann's (2008) unobserved heterogeneous, model-based, clustering approach. The two approaches are similar in that both model the hyperprior as an unknown. But, they differ in that Frühwirth-Schnatter & Kaufmann (2008) models the hyperprior with a finite mixture, whereas our hyperprior is almost surely an infinite ordered mixture distribution.<sup>4</sup>

Formally, our nonparametric hierarchical prior for the panel of skill measurements,  $\alpha_{im}$ , is defined as

$$\alpha_{im}|a_{\alpha,im}, h_{\alpha,im}^2 \sim N(a_{\alpha,im}, h_{\alpha,im}^2), \quad (4)$$

$$(a_{\alpha,im}, h_{\alpha,im}^2) \sim G_\alpha, \quad (5)$$

---

<sup>4</sup>Chan & Koop (2014) also use Frühwirth-Schnatter & Kaufmann's (2008) unobserved heterogeneous, model-based, clustering approach to group together structural break series.

where  $G_\alpha$  is a unknown hyperprior distribution. A similar conditional prior representation exists for our nonparametric hierarchical prior of the multiple-change-point variances,  $\sigma_{im}^2$ , where

$$\sigma_{im}^2 | h_{\sigma,im}^2 \sim \text{Inv-Gamma}(\nu_\sigma/2, \nu_\sigma h_{\sigma,im}^2/2), \quad (6)$$

$$h_{\sigma,im}^2 \sim G_\sigma. \quad (7)$$

Note the shape hyperparameter,  $\nu_\sigma$ , is assumed to be known and is equal to the same value for all  $J$  mutual funds. In Eq. (6) and (7),  $h_{\sigma,im}^2$  is the unknown scale hyperparameter whose distribution,  $G_\sigma$ , is unknown.<sup>5</sup>

We define the nonparametric hierarchical prior for the transition probabilities,  $q_i$ , to be

$$q_i | j_i, k_i \sim \text{Beta}(j_i, k_i - j_i + 1), \quad (8)$$

$$(j_i, k_i) \sim G_q, \quad (9)$$

where  $k_i \in \mathbb{N}$ ,  $j_i \in \{1, \dots, k_i\}$ , and  $G_q$  is the unknown hyperprior distribution. Unlike the regime parameters' hierarchical priors, which model how the parameters are distributed across both fund and their regimes, Eq. (8)-(9) only models how the  $q_i$ s are distributed across the  $J$  funds.<sup>6</sup>

Mixing the beta distribution in Eq. (8) over its hyperparameters leads to the flexible Bernstein prior. Research by Petrone & Wasserman (2002) shows the posterior of the Bernstein prior converging to the true data generating density of a continuous and bounded distribution residing on the unit interval. To our knowledge a hierarchical beta distribution has never been proposed before as a prior for an unknown parameter defined on the unit interval (see Fisher (2017)).

Hierarchical priors like those above are desirable since they facilitate the sharing of information across the cross-section of funds and their regimes through the priors' unknown hyperparameters,  $a_{\alpha,im}$ ,  $h_{\alpha,im}^2$ ,  $h_{\sigma,im}^2$ ,  $j_i$  and  $k_i$ . For instance, the observed performance of a highly skilled manager or a particular trading strategy gets reflected in the posterior of the performance prior's hyperparameters  $a_{\alpha,im}$  and  $h_{\alpha,im}^2$ . These posterior hyperparameter estimates then help reduce the uncertainty around how another fund might perform if it were to hire away the manager or adopt the same trading strategy. This sharing of information

---

<sup>5</sup>Because neither the conditional variances, nor their log transform, are normally distributed, our multiple-change-point panel model does not have the mixture innovation representation of Giordani & Kohn (2008). Therefore, we cannot sample the change point process with Gerlach et al.'s (2000) fast and efficient algorithm.

<sup>6</sup>If in the future we were interested in investigating the persistence of a particular regime,  $m_i$ , for instance, the skill of a highly successful manager, there is nothing in our approach that would prohibit it from being extended to the regime specific transition probability case.

through the learning of the hyperparameters is not possible when the parameters come from a prior distribution whose hyperparameters are fixed apriori.

Parametric hierarchical priors where the hyperprior distribution is assumed to be known have been used before to infer the parameters of change points and structural break models (see Pesaran et al. (2006), Koop & Potter (2007), Geweke & Jiang (2011) and Maheu & Song (2014)). In the context of mutual fund performance, assuming the hyperprior is known amounts to making a strong assumption about how the change point parameters are distributed across the funds and their regimes. For instance, assuming a normal, inverse-gamma hyperprior for  $G_\alpha$  explicitly leads to the alphas being distributed over the population of mutual funds and their regimes in a unimodal fashion whose mode is the global average skill level. Given the abundant information contained in the cross-section of a longitudinal data set, this parametric assumption causes the posterior estimates to globally shrink towards the average fund-regime level. Alphas for the highly skilled and unskilled fund-regimes will end up looking ordinary like that of the skill level for a typical fund-regime. To overcome this global shrinkage we let the hyperpriors be unknown and, like the hyperparameters, learn about the hyperprior distributions by modeling them with DP priors.

### 3.1 Dirichlet process prior

The Dirichlet process has received considerable attention from the Bayesian nonparametric community as a prior to the mixture weights and location parameters of a infinite mixture representation of a unknown distribution (see Escobar & West (1995) and Hjort et al. (2010), and the chapters therein, for a good introduction to the DP). DPs have also been used extensively in economics and finance to model unknown distributions (see Chib & Tiwari (1988), Chib & Hamilton (2002), Hirano (2002), Jensen & Maheu (2010), and Bassetti et al. (2014)). In addition, Song (2014), Dufays (2015), Jochmann (2015) and Jin & Maheu (2016) apply the DP to structural break models to allow for an infinite number of possible regimes.

One reason for this attention is that the DP delivers a sparse mixture representation of unknown distributions. It is this sparsity that leads to the local shrinkage in the regime parameter estimates. Another reason is the DPs ease of use. As a conjugate distribution the DP lends itself to a straightforward sampler, one that quickly converges to draws from the posterior distribution of the nonparametric hierarchical prior. Under the DP prior our nonparametric hierarchical prior also nests the parametric hierarchical priors used by Chib (1998), Pesaran et al. (2006) and Geweke & Jiang (2011). These properties lead us to apply

the DP prior to the three hierarchical priors defined in Eq. (4)–(9).

Let the unknown hyperpriors have the following DP priors

$$G_\alpha \sim DP(\eta_\alpha, G_{\alpha,0}), \quad (10)$$

$$G_\sigma \sim DP(\eta_\sigma, G_{\sigma,0}), \quad (11)$$

$$G_q \sim DP(\eta_q, G_{q,0}). \quad (12)$$

The non-negative scalars,  $\eta_\alpha$ ,  $\eta_\sigma$ , and  $\eta_q$ , are the concentration parameters of the DPs. Their value is important since it determines the propensity of clustering and the variance in the DP; e.g., for the hyperprior  $G_\sigma$ ,  $\text{Var}[G_\sigma(A)] = [G_{\sigma,0}(A)(1 - G_{\sigma,0}(A))]/(1 + \eta_\sigma)$ , where  $A$  is a measurable subset of  $h_{\sigma,im}^2$  domain. A larger concentration parameter causes the DP to be less parsimonious and to have a lower propensity to cluster the hyperparameters together. For example, when  $\eta_\sigma \rightarrow \infty$ , each  $\sigma_{im}^2$  belongs to its own unique group.

To simplify our exposition we initially treat the concentration parameters as known values, but, empirically infer their value by sampling from their conditional posterior distribution. When sampled we assume a Log-Logistic(1, 1) distribution for their prior.<sup>7</sup> Under this prior the probability of the DP having a second cluster,  $u_\bullet \equiv \eta_\bullet/(1 + \eta_\bullet)$ , is uniform over the unit interval. This probability goes to zero as  $\eta_\bullet \rightarrow 0$ , whereas it goes to one as  $\eta_\bullet \rightarrow \infty$ . We can then address the amount of empirical evidence in favor of a parametric hierarchical prior by measuring the posterior density of  $u_\bullet$  at zero.

$G_{\sigma,0}$ ,  $G_{\alpha,0}$ , and  $G_{q,0}$  are the base distributions of the DP priors and must be specified by the econometrician. These distributions equal the expected distribution of the DP prior; e.g.,  $G_{\sigma,0} \equiv E[G_\sigma]$ . In our nonparametric hierarchical priors, the base distribution represents our initial view about how the hyperparameters are distributed over the relevant population. In this paper we choose base distributions that bring as little prior information about the population distribution to the analysis as possible.<sup>8</sup>

In our empirical application we set  $G_{\alpha,0}$  equal to the conditional conjugate, normal, inverse-gamma distribution,  $\text{NIG}(a_\alpha, h_\alpha^2 | a_0, h_\alpha^2/\kappa_0, \nu_0/2, \nu_0 h_0^2/2)$ , where the mean is  $a_0 = 0$ , the scaling of the variance is  $\kappa_0 = 0.1$ , the shape is  $\nu_0/2 = 0.005$ , and the scale is  $\nu_0 h_0^2/2 = 0.005 \times 0.01$ . With the NIG base distribution it follows from Eq. (4) that our initial guess for how the alphas are distributed over the funds and their regimes is represented by the

---

<sup>7</sup>The probability density function for the Log-Logistic(1, 1) prior distribution of  $\eta_\bullet$  is  $1/(1 + \eta_\bullet)^2$  when  $\eta_\bullet > 0$ , zero otherwise.

<sup>8</sup>In subsequent research we have found that the empirical results in Section 5 are robust to the choice of the base distribution, as long as there is sufficient variation in the prior distributions drawn from the nonparametric hierarchical prior.

prior predictive Student-t distribution

$$\begin{aligned}
E_{G_{\alpha,0}}[E_{G_{\alpha}}[\pi_{\alpha}(\alpha|G_{\alpha})]] &= \int N(\alpha|a_{\alpha}, h_{\alpha}^2) NIG(a_{\alpha}, h_{\alpha}^2 | a_0, h_{\alpha}^2/\kappa_0, \nu_0/2, \nu_0 h_0^2/2) d(a_{\alpha}, h_{\alpha}^2), \\
&= t_{\nu_0} \left( \alpha \middle| a_0, \left( \frac{\kappa_0 + 1}{\kappa_0} \right) \nu_0 h_0^2 \right), \tag{13}
\end{aligned}$$

where the first expectation is taken with respect to the unknown hyperprior,  $G_{\alpha}$ , which equals  $G_{\alpha,0}$ .

Other than the efficient market hypothesis argument that the average level of skill over the population is zero ( $a_0 = 0$ ), there is little theory to guide us as to what the population distribution for the alphas should look like. Instead, we intentionally set the degrees of freedom,  $\nu_0$ , to be less than two in order for the variance of the prior predictive distribution of the alphas to be undefined. With an undefined variance our initial guess at the distribution of skill is very diffuse and uninformative.<sup>9</sup>

For the prior of the variances we let  $G_{\sigma,0} \equiv \text{Gamma}(h_{\sigma}^2 | c_0, 1/b_0)$ , with shape,  $c_0 = 0.3$ , and scale,  $b_0 = 0.001$ . We also set the scale of the conditional, inverse-gamma, prior in Eq. (6) equal to  $\nu_{\sigma} = 0.5$ . Dubey (1970) shows that a gamma hyperprior to the conditional inverse-gamma distribution results in the proper beta prime, prior predictive, distribution

$$E_{G_{\sigma,0}}[E_{G_{\sigma}}[\pi_{\sigma}(\sigma^2|G_{\sigma})]] = \frac{\left(\frac{\sigma^2}{w} + 1\right)^{c_0 + \nu_{\sigma}/2} \left(\frac{\sigma^2}{w}\right)^{c_0 - 1}}{w \text{Beta}(c_0, \nu_{\sigma}/2)}, \tag{14}$$

where  $w = \nu_{\sigma}/(2b_0)$  and  $\text{Beta}(c_0, \nu_{\sigma}/2)$  is the beta function.

For the chosen values of  $c_0$ ,  $b_0$  and  $\nu_{\sigma}$ , the expected value of the proper beta prime is undefined. However, its median is 402.99 and its lower and upper quantile are 26.5 and 7966.6, respectively. Hence, our initial guess at how the variances are distributed over the funds and their regimes is very diffuse. This ensures that the posterior predictive distribution of the variances will be based on the longitudinal return data.

In the hierarchical prior for the  $q_i$ s, Eq. (8) is equivalent to the distribution of a  $j_i$ th order statistic for a random sample of  $k_i$  uniform draws. By the adding-up property of the Bernstein polynomial, the prior predictive for the  $q_i$ s is thus a uniform distribution on the unit interval

$$E_{G_{q,0}}[E_{G_q}[\pi_q(q|G_q)]] = \text{Uniform}(q|0, 1), \tag{15}$$

---

<sup>9</sup>For our nonparametric hierarchical prior to have a proper prior predictive distribution the DP's base distribution must be a proper distribution. As a result we cannot use a Jeffereys prior for the base distribution.

regardless of the value of  $k_i$ . The base distribution,  $G_{q,0}$ , is characterized by

$$j_i | k_i \sim \text{Uniform}(1, \dots, k_i), \quad (16)$$

$$k_i - 1 \sim \text{Geometric}(\xi_0). \quad (17)$$

Even though the value of  $k_i$  does not affect the prior predictive uniform distribution, it does affect the variation in the random realizations from the nonparametric hierarchical prior. Small values of  $k_i$  (a wider beta kernel) restrict the variation in the random realizations, whereas a large  $k_i$  (a tighter beta kernel) enhances the variation.

Shrinkage is also controlled by the value of  $k_i$ . When  $k_i = 1$  there is no local shrinkage in a cluster's  $q_i$ s, whereas, when  $k_i \rightarrow \infty$ , there is complete local shrinkage. Given these properties we desire more weight be placed on larger values of  $k_i$ , so we set  $\xi_0 = 1/200$  such that  $Pr(k_i = 1) = 1/200$  and  $E[k_i] = 200$ .

Besides using these specific base distributions for the nonparametric hierarchical priors in our empirical analysis we also use them as the known hyperpriors of the parametric hierarchical priors. The prior predictives are then the same for both the parametric and nonparametric hierarchical priors.

### 3.2 Posterior DP

After the returns from the panel of mutual funds are observed the posterior of the hierarchical priors are calculated by updating the base distributions of the hyperprior's DP. Since we are interested in the behavior of skill it would be natural to discuss the posterior distribution of the hyperparameters  $a_{\alpha,im}$  and  $h_{\alpha,im}^2$ . However, their posterior has more moving parts and is more complicated than the description of  $h_{\sigma,im}^2$  posterior so we only explain the posterior in terms of the variance's hyperparameters (see the online Appendix for complete details on all the posteriors).

Suppose it were possible to observe the hyperparameters of all  $J$  funds and their  $\mathcal{N}_A = \sum_{i=1}^J s_{i,T_i}$  total in-sample regimes.<sup>10</sup> In terms of the hyperparameters for the prior of the variances let  $\mathcal{H}_{\sigma,S_T} = \{\mathbf{h}_{\sigma,is_{T_i}}\}_{i=1}^J$ , where  $\mathbf{h}_{\sigma,is_{T_i}} = (h_{\sigma,i1}^2, \dots, h_{\sigma,is_{T_i}}^2)'$  contains the  $\mathcal{N}_A$  scale hyperparameters from all the funds and their in-sample regimes. If we could actually observe  $\mathcal{H}_{\sigma,S_T}$  it would contain  $\mathcal{N}_A$  independent realizations from the unknown hyperprior,  $G_{\sigma}$ .

---

<sup>10</sup>In practice the hyperparameters are never observed, hence, their uncertainty will need to be integrated away after the information found in the return data of the funds and their regimes is used to quantify their posterior distribution.

Combining the hypothetical hyperparameter data in  $\mathcal{H}_{\sigma, S_T}$  with the conjugate DP prior for  $G_\sigma$  results in the posterior DP hyperprior

$$G_\sigma | \mathcal{H}_{\sigma, S_T} \sim DP \left( \eta_\sigma + \mathcal{N}_A, \frac{\eta_\sigma G_{\sigma,0} + \sum_{i=1}^J \sum_{m=1}^{s_{i,T_i}} \delta_{h_{\sigma,im}^2}}{\eta_\sigma + \mathcal{N}_A} \right), \quad (18)$$

where  $\delta_{h_{\sigma,im}^2}$  is a degenerative distribution at the location  $h_{\sigma,im}^2$ . Since the concentration parameter of the posterior DP equals  $\eta_\sigma + \mathcal{N}_A$ , the uncertainty around  $G_\sigma$  has declined. The base distribution of the posterior DP distribution equals

$$\frac{\eta_\sigma G_{\sigma,0} + \sum_{i=1}^J \sum_{m=1}^{s_{i,T_i}} \delta_{h_{\sigma,im}^2}}{\eta_\sigma + \mathcal{N}_A},$$

which illustrates how the initial guess,  $G_{\sigma,0}$ , has been updated with the empirical distribution of the “observed” hyperparameters,  $M_{S_T}^{-1} \sum_{i=1}^J \sum_{m=1}^{s_{i,T_i}} \delta_{h_{\sigma,im}^2}$ .

From the base distribution of the posterior DP we can see that as long as  $\eta_\sigma$  is finite every occurrence of the hyperparameters,  $h_{\sigma,im}^2$ , adds more empirical information to the guess of  $G_\sigma$ , while giving less weight to the initial guess,  $G_{\sigma,0}$ . If the DP prior’s concentration parameter were infinite ( $u_\bullet = 1$ ), the variance of the DP would be zero and one would learn nothing from the observed data.

At the other end of the spectrum is  $\eta_\bullet = 0$  ( $u_\bullet = 0$ ). In this situation every realization of the hyperparameter has the exact same value equal to a single draw from  $G_{\sigma,0}$ . When this is the case the nonparametric hierarchical prior is equivalent to the parametric hierarchical prior having the hyperprior  $G_{\sigma,0}$ .

### 3.3 Clustering

According to the base distribution of the posterior DP in Eq. (18), a realization from the posterior is a draw from either the initial base distribution,  $G_{\sigma,0}$ , or from the discrete collection of  $\mathcal{N}_A$  degenerative distributions,  $\delta_{h_{\sigma,im}^2}$ . As explained in Section 3.2 drawing from the discrete distribution depends on the value of  $\eta_\sigma$ . These draws create ties between the realizations, and, hence, creates  $\mathcal{K}_\sigma$  clusters consisting of the unique hyperparameters,  $h_{\sigma,c}^{2*}$ ,  $c = 1, \dots, \mathcal{K}_\sigma$ , where  $\mathcal{K}_\sigma \leq \mathcal{N}_A$ .<sup>11</sup>

The posterior DP in Eq. (18) can be rewritten in terms of the clusters as

$$G_\sigma | H_\sigma^*, \mathbf{n}_\sigma \sim DP \left( \eta_\sigma + \mathcal{N}_A, \frac{\eta_\sigma G_{\sigma,0} + \sum_{c=1}^{\mathcal{K}_\sigma} n_{\sigma,c} \delta_{h_{\sigma,c}^{2*}}}{\eta_\sigma + \mathcal{N}_A} \right), \quad (19)$$

---

<sup>11</sup>This clustering is similar to the model-based clustering of Frühwirth-Schnatter & Kaufmann (2008) except the clustering in our DP is in the hyperparameters and the number of clusters,  $\mathcal{K}_\sigma$ , is unknown.

where  $H_\sigma^* = (h_{\sigma,1}^{2*}, \dots, h_{\sigma,\mathcal{K}_\sigma}^{2*})'$  is a vector comprised of the  $\mathcal{K}_\sigma$ -uniquely valued  $h_{\sigma,im}^2$ s, and  $\mathbf{n}_\sigma = (n_{\sigma,1}, \dots, n_{\sigma,\mathcal{K}_\sigma})'$ , where  $n_{\sigma,c}$  is the number of  $h_{\sigma,im}^2$ s equal to  $h_{\sigma,c}^{2*}$ . It follows that  $\sum_{c=1}^{\mathcal{K}_\sigma} n_{\sigma,c} = \mathcal{N}_\mathcal{A}$ .

If the membership to the  $c$ th cluster includes a large number of fund-regimes then  $n_{\sigma,c}$  will be large relative to  $\mathcal{N}_\mathcal{A}$ . Then according to Eq. (19), the  $c$ th cluster will have a greater chance of a new fund-regime being assigned to it. A nice manageable number of groups occurs, and a sparse, parsimonious, mixture representation of the nonparametric prior is found in

$$\begin{aligned} E_{G_{\sigma,0}}[E_{G_\sigma}[\pi_\sigma(\sigma^2|G_\sigma, H_\sigma^*, \mathbf{n}_\sigma)]] &= \frac{\eta_\sigma}{\eta_\sigma + \mathcal{N}_\mathcal{A}} \int IG(\sigma^2|\nu_\sigma/2, \nu_\sigma h_\sigma^{2*}/2) dG_{\sigma,0}(h_\sigma^{2*}) \\ &+ \sum_{c=1}^{\mathcal{K}_\sigma} \frac{n_{\sigma,c}}{\eta_\sigma + \mathcal{N}_\mathcal{A}} IG(\sigma^2|\nu_\sigma/2, \nu_\sigma h_{\sigma,c}^{2*}/2). \end{aligned} \quad (20)$$

By partitioning the funds and their regimes into large, small, and possibly singleton ( $n_{\sigma,c} = 1$ ) groups where the members have similarly distributed fund-regime parameters, the DP hyperprior is robust to global shrinkage. Highly skilled funds and extraordinary performing trading strategies are able to speak for themselves in determining their group. Members then share in the information found in their group and the regime parameters shrink locally towards the group's hyperparameters.<sup>12</sup>

Going forward we drop the expectation operators from the predictive distributions knowing that the prior distributions are unknown and estimated with the expected value of their posterior. For instance, we denote the posterior predictive distribution of the variances as

$$\pi_\sigma(\sigma^2|\mathcal{Y}) = \int E_{G_{\sigma,0}}[E_{G_\sigma}[\pi_\sigma(\sigma^2|G_\sigma, \eta_\sigma, H_\sigma^*, \mathbf{n}_\sigma)]] \pi(\eta_\sigma, H_\sigma^*, \mathbf{n}_\sigma|\mathcal{Y}) d\eta_\sigma dH_\sigma^* d\mathbf{n}_\sigma.$$

## 4 Posterior simulation

In this section we describe our approach to sampling from the posterior distribution. We divide the unknown change point processes, the regime parameters, and the unknown prior distributions into natural blocks and sample each block from its conditional posterior distribution. For a thorough and complete description of the sampler and the posterior distributions please refer to the online Appendix.

<sup>12</sup>The sampler in Maheu & Song (2014) requires the hierarchical prior to be the joint conjugate prior  $\pi(\mu, \sigma^2|a_\mu, h_\mu^2, h_\sigma^2)$ , where  $(a_\mu, h_\mu^2, h_\sigma^2) \sim G_{\mu,\sigma}$ . Because of the Bayesian preference for sparsity such multidimensional DP priors for the hyperprior may have fewer posterior clusters, and hence, possibly be less flexible than our priors.

There are two types of blocks: within-fund blocks and across-fund blocks. In the joint likelihood of Eq. (3) we see the within-fund blocks are conditionally independent across funds and consequently can be parallelized. For a given fund, these blocks include drawing the change point process, the transition probability, and the fund-regime parameters. The across-fund parameters are associated with the nonparametric hierarchical priors. They involve drawing the hyperparameters and concentration parameters. These across-fund blocks are conditionally independent across hyperparameters.

#### 4.1 Sampling the within-fund unknowns

Draws of the in-sample, change point process,  $S_{i,T_i}$ , from the conditional posterior distribution,  $\pi(S_{i,T_i}|Y_{i,T_i}, \alpha_i, \sigma_i^2, q_i)$ , are made with the forward-filter, backward-smoother, sampler of Chib (1996). As we have already pointed out the draws of  $S_{i,T_i}$  are parallelized over the  $J$  funds. Note that these latent change point draws only go to the end of the return series,  $Y_{i,T_i}$ . To sample the out-of-sample change points we apply the same forward-backward approach up to the  $M_i$ th regime.

Given the draw of  $S_{i,T_i}$ , and the hyperparameters,  $j_i$  and  $k_i$ , we then draw  $q_i$  from

$$\pi(q_i|\mathcal{T}_i, s_{i,T_i}, j_i, k_i) = \text{Beta}(q_i | j_i + s_{i,T_i} - 1, k_i - j_i + 1 + (\mathcal{T}_i - s_{i,T_i})), \quad (21)$$

where  $s_{i,T_i} - 1$  is the number of change-points out of  $\mathcal{T}_i - 1$  independent Binomial trials (the minus one accounts for  $s_{i,T_i} = 1$  with probability one).

To draw the elements of  $\alpha_i$  and  $\sigma_i^2$  we condition on the current draw of  $S_{i,T_i}$  and the hyperparameters,  $(a_{\alpha,im}, h_{\alpha,im}^2)$  and  $h_{\sigma,im}^2$ ,  $m = 1, \dots, M_i$ . Given the conjugate hierarchical prior for the variances found in Eq. (6),  $\sigma_{im}^2$  is drawn from the conditional posterior by sampling from the marginal distribution

$$\pi(\sigma_{im}^2 | \{y_{i,t} : s_{i,t} = m\}, \alpha_{im}, h_{\sigma,im}^2) \propto \prod_{t:s_{i,t}=m} N(y_{i,t} | \alpha_{im}, \sigma_{im}^2) \text{IG}(\sigma_{im}^2 | \nu_\sigma/2, \nu_\sigma h_{\sigma,im}^2).$$

The  $\alpha_{im}$ s are then drawn from the conditional posteriors

$$\pi(\alpha_{im} | \{y_{i,t} : s_{i,t} = m\}, \sigma_{im}^2, a_{\alpha,im}, h_{\alpha,im}^2) \propto \prod_{t:s_{i,t}=m} N(y_{i,t} | \alpha_{im}, \sigma_{im}^2) N(\alpha_{im} | a_{\alpha,im}, h_{\alpha,im}^2).$$

Note that when the  $i$ th fund's regime,  $m$ , is greater than  $s_{i,T_i}$  the posteriors for  $\alpha_{im}$  and  $\sigma_{im}^2$  have no return data to condition on. To draw the regime parameters for these  $M_i - s_{i,T_i}$  out-of-sample regimes we sample from the posterior conditional priors

$$\sigma_{im}^2 \sim \text{IG}(\sigma_{im}^2 | \nu_\sigma/2, \nu_\sigma h_{\sigma,im}^2), \quad (22)$$

$$\alpha_{im} \sim N(y_{i,t} | \alpha_{im}, \sigma_{im}^2) N(\alpha_{im} | a_{\alpha,im}, h_{\alpha,im}^2), \quad (23)$$

where the hyperparameters,  $h_{\sigma,im}^2$ ,  $a_{\alpha,im}$ , and  $h_{\alpha,im}^2$ , for  $m > s_{i,T_i}$ , are draws from the out-of-sample posterior hyperprior described in Section 4.2.2.

## 4.2 Sampling the across-fund unknowns

Given each fund's posterior draw of  $\alpha_i$ ,  $\sigma_i^2$ , and  $q_i$ , the sampler moves on to drawing the across-fund unknowns associated with the nonparametric hierarchical representation of  $\pi_\alpha$ ,  $\pi_\sigma$ , and  $\pi_q$ . All three priors have hierarchical conjugate hyperpriors for the DP priors and their base distributions. This allows the hyperparameters to be quickly and efficiently drawn with the two-step DP sampler of West et al. (1994), MacEachern & Müller (1998), and described by Neal (2000) in his Algorithm 2.

The DP two-step sampler first sequentially assigns every regime of every fund a random hyperparameter and in the process partitions the fund-regimes into groups having the same hyperparameters. In the second step the hyperparameters of each regime group are independently drawn.

### 4.2.1 Transition probability hyperparameters

For the prior of the transition probabilities, the first step of the DP sampler sequentially draws the  $j_i$  and  $k_i$  along with the unknown cluster assignment variable,  $z_{q,i}$ , for  $i = 1, \dots, J$ . The assignment variable,  $z_{q,i} = c$ , where  $c = 1, \dots, \mathcal{K}_q$ , when  $j_i = j_c^*$  and  $k_i = k_c^*$ , and  $j_c^*$ , and  $k_c^*$  are the unique hyperparameters of the  $c$ th cluster. These draws and assignments are all done by applying the Chinese restaurant process sampling algorithm to the  $j_i$ s and  $k_i$ s (see Teh (2010) for details on the Chinese restaurant process).

In the second step, separate draws are made from the posterior  $\pi(j_c^*, k_c^* | \{q_i : z_{q,i} = c\})$ , for  $c = 1, \dots, \mathcal{K}_q$ . Since this conditional posterior distribution is not an analytical distribution, we design an efficient Metropolis-Hasting (MH) sampling scheme to draw from this distribution. The technical details of this MH sampler are described in the online Appendix.

### 4.2.2 Regime hyperparameters

To sample the hyperparameters  $a_{\alpha,im}$ ,  $h_{\alpha,im}^2$  and  $h_{\sigma,im}^2$ , we partition the hyperparameters, the alphas,  $\alpha_{im}$ , and the variances,  $\sigma_{im}^2$ , into in-sample ( $m \leq s_{i,T_i}$ ) and out-of-sample ( $s_{i,T_i} < m \leq M_i$ ) regimes. Because  $M_i$  exceeds  $s_{i,T_i}$ , our sampler always has out-of-sample unknowns where there is no return data.<sup>13</sup> As a result the out-of-sample hyperparameter

<sup>13</sup>If at any point in the sampler the inequality,  $s_{i,T_i} < M_i$ , did not hold we increased  $M_i$  and started the entire sampler over.

draws are realizations from the hyperprior's posterior predictive distribution.

Let the in-sample hyperparameters,  $a_{\alpha,im}$ ,  $h_{\alpha,im}^2$ , and,  $h_{\sigma,im}^2$ , where  $m = 1, \dots, s_{i,T_i}$ , and  $i = 1, \dots, J$ , be stored in  $\mathcal{A}_{\alpha,s_T}$ ,  $\mathcal{H}_{\alpha,s_T}$ , and  $\mathcal{H}_{\sigma,s_T}$ , respectively. Each set contains the  $N_{\mathcal{A}}$  in-sample hyperparameters. Also let the  $N_{\mathcal{A}}$ -element set of in-sample alphas and variances be defined as  $\mathbf{A}_{s_T} = \{\alpha_{im} : m = 1, \dots, s_{i,T_i}, i = 1, \dots, J\}$  and  $\mathbf{\Sigma}_{s_T} = \{\sigma_{im}^2 : m = 1, \dots, s_{i,T_i}, i = 1, \dots, J\}$ , respectively.

The two base distributions

$$G_{\alpha,0} \equiv \text{NIG}(a_{\alpha}, h_{\alpha}^2 | a_0, h_{\alpha}^2 / \kappa_0, \nu_0 / 2, \nu_0 h_0^2 / 2), \quad G_{\sigma,0} \equiv \text{Gamma}(h_{\sigma}^2 | c_0, 1/b_0),$$

are conjugate to the normally distributed conditional prior in Eq. (4) and the inverse-gamma distributed conditional prior in Eq. (6), respectively. Hence, drawing the in-sample hyperparameters from  $\pi(\mathcal{A}_{\alpha,s_T}, \mathcal{H}_{\alpha,s_T} | \mathbf{A}_{s_T}, \eta_{\alpha})$  and  $\pi(\mathcal{H}_{\sigma,s_T} | \mathbf{\Sigma}_{s_T}, \eta_{\sigma})$  is a straight-forward application of the two-step DP sampler (see the Internet Appendix for the technical details of these two-step draws).

Since the out-of-sample hyperparameters are exchangeable the order in which they are drawn does not matter. So for a tractable description of the out-of-sample draws one can think of the sampler starting with the first fund and then sequentially sampling the first fund's out-of-sample hyperparameters for the regimes,  $s_{1,T_1} + 1, \dots, M_1$ , before moving on to drawing the out-of-sample hyperparameters for the second fund. The sampler then continues sequentially drawing the out-of-sample hyperparameters, fund by fund, until the hyperparameters have been drawn for all  $J$  mutual funds.

Formally, the draw of the  $i$ th fund's  $m$ th out-of-sample hyperparameters,  $(a_{\alpha,im}, h_{\alpha,im}^2)$ , where  $m > s_{i,T_i}$ , are sampled from the conditional posterior predictive distribution

$$\frac{1}{\eta_{\alpha} + M_{s_T} + M_{im}} \left[ \eta_{\alpha} \text{NIG}(a_{\alpha,im}, h_{\alpha,im}^2 | a_0, h_{\alpha}^2 / \kappa_0, \nu_0 / 2, \nu_0 h_0^2 / 2) + \sum_{i' < i} \sum_{m'=1}^{M_{i'}} \delta_{(a_{\alpha,i'm'}, h_{\alpha,i'm'}^2)} + \sum_{l=1}^{m-1} \delta_{(a_{\alpha,il}, h_{\alpha,il}^2)} \right], \quad (24)$$

where  $M_{im}$  is the number of out-of-sample hyperparameters drawn before sampling  $(a_{\alpha,im}, h_{\alpha,im}^2)$ .<sup>14</sup> Posterior draws of the out-of-sample hyperparameters,  $h_{\sigma,im}^2$ , where  $m > s_{i,T_i}$ , are similarly made fund by fund (see the online Appendix for the exact form of the posterior distribution).

<sup>14</sup>Note that the out-of-sample hyperparameters will cluster into groups having the same unique values,  $(a_{\alpha,c}^*, h_{\alpha,c}^{2*})$ ,  $c = 1, \dots, \mathcal{K}_{\alpha}$ . Therefore, we could have written Eq. (24) in terms of the unique hyperparameter values, however, this would have required additional notation and a more involved explanation.

### 4.3 Concentration parameter draws

Posterior draws of the concentration parameters,  $\eta_\mu$ ,  $\eta_\sigma$  and  $\eta_q$ , are all made using the same MH algorithm. The density for the MH proposal draw  $\eta'_\bullet$  is  $\eta'_\bullet = \eta_\bullet u_\bullet / (1 - u_\bullet)$  where  $u_\bullet \sim \text{Uniform}(0, 1)$ . Given the Log-Logistic prior for the concentration parameter it follows that  $\pi(\eta'_\bullet | \eta_\bullet) = \eta_\bullet / (\eta_\bullet + \eta'_\bullet)^2$ . Step-by-step details of the MH sampler are found in the online Appendix.

### 4.4 Posterior predictive

To compute the posterior predictive distributions we integrate away the uncertainty found in the hyperparameters of the hierarchical prior and the DP's concentration parameters by numerically averaging the conditional distribution of the parameter over the posterior draws of these unknowns. For example, the predictive distribution for the variances is calculated as

$$\begin{aligned} \pi_\sigma(\sigma^2 | \mathcal{Y}) \approx & R^{-1} \sum_{r=1}^R \left[ \frac{\eta_\sigma^{(r)}}{\eta_\sigma^{(r)} + M_{S_T}^{(r)}} f_{PBP}(\sigma^2 | c_0, \nu_\sigma/2, 1, w) \right. \\ & \left. + \sum_{c=1}^{\mathfrak{K}_\sigma^{(r)}} \frac{n_{\sigma,c}^{(r)}}{\eta_\sigma^{(r)} + M_{S_T}^{(r)}} f_{IG} \left( \sigma^2 \mid \nu_\sigma/2, \nu_\sigma h_{\sigma,c}^{2*(r)}/2 \right) \right], \end{aligned} \quad (25)$$

which averages the conditional posterior predictive distribution from Eq. (20) over the  $R$  posterior draws  $h_{\sigma,c}^{2*(r)}$ ,  $\eta_\sigma^{(r)}$  and  $M_{S_T}^{(r)}$ ,  $r = 1, \dots, R$ . Similar calculations are performed for the posterior predictive densities,  $\pi_\alpha(\alpha | \mathcal{Y})$ , and  $\pi_q(q | \mathcal{Y})$  (see the online Appendix for these formulas).

## 5 Mutual fund application

An important question mutual fund investors ask is, do above market returns by a mutual fund today lead to above market returns in the future? Along with an answer to this question investors would also like to know if the performance of a fund is linked to the success of its manager. These and other performance related questions have been addressed in the academic and profession literature dating back to Jensen (1968).<sup>15</sup> Bollen & Whaley (2009) apply a single-change-point model to a cross-section of hedge funds, but, to our knowledge a multiple-change-point panel model has never been used to measure mutual fund skill and its persistence.

---

<sup>15</sup>See Elton & Gruber (2013) for a review of the mutual fund performance literature.

To measure skill and its persistence we analyze mutual fund performance with the multiple-change-point, four-factor-risk, model

$$y_{i,t} = \alpha_{i,s_{i,t}} + \beta_{R,i,s_{i,t}} \text{RMRF}_t + \beta_{S,i,s_{i,t}} \text{SMB}_t + \beta_{H,i,s_{i,t}} \text{HML}_t + \beta_{M,i,s_{i,t}} \text{MOM}_t + \sigma_{i,s_{i,t}} \epsilon_{i,t}, \quad (26)$$

where  $y_{i,t}$  is the risk-free adjusted, gross rate of return generated by the  $i$ th mutual fund in month  $t$ . The risk factors are Fama & French's (1993) three factors of excess market returns,  $\text{RMRF}_t$ , market size,  $\text{SMB}_t$ , and book-to-market,  $\text{HML}_t$ , plus Carhart's (1997) momentum portfolio factor,  $\text{MOM}_t$ . The intercept,  $\alpha_{i,s_{i,t}}$ , is a change point version of the Jensen (1968) alpha that allows a fund's level of performance to change with its change point process,  $s_{i,t}$ .

We assume the prior for  $\alpha_{i,s_{i,t}}$  and the risk factor parameters,  $\beta_{R,i,s_{i,t}}$ ,  $\beta_{S,i,s_{i,t}}$ ,  $\beta_{H,i,s_{i,t}}$ , and  $\beta_{M,i,s_{i,t}}$ , consists of the independent marginals denoted by  $\pi_\alpha$ ,  $\pi_{\beta_R}$ ,  $\pi_{\beta_S}$ ,  $\pi_{\beta_H}$ , and  $\pi_{\beta_M}$ , respectively. Each prior is unknown and modeled with the nonparametric hierarchical prior defined in Eq. (4), (5), and (10). In this paper we only discuss our findings for the distributions of the alphas but in future work we plan to report and analyze the distributions of the betas.

## 5.1 Longitudinal mutual fund data

Our longitudinal data set consists of the monthly mutual fund returns investigated by Jones & Shanken (2005).<sup>16</sup> For ease of comparison we annualized the returns by multiplying the monthly return by twelve so the alphas are yearly above market rates of return. The data is comprised of the returns from every US domestic, equity, mutual fund that existed between January 1961 to June 2001.<sup>17</sup> Each fund has at least a years worth of return data and on average the funds have 77.3 months of returns. Survivorship bias is not a issue since the panel includes the returns from all 1,293 funds that did not survive to the end of the sample. In total the data consists of a unbalanced panel of 396,820 monthly observations across 5,136 domestic equity funds.

## 5.2 Posterior inference

Posterior inference is made using both nonparametric and parametric hierarchical priors. As we explained in Section 3.1, neither approach is at an initial disadvantage since we set the base distribution of the DP hyperpriors equal to the parametric hyperpriors. Both approaches have similar samplers but the parametric model only needs to draw a single

---

<sup>16</sup>We would like to thank Chris Jones for providing us with their data.

<sup>17</sup>Funds were excluded if they made substantial investments in asset classes other than domestic equities.

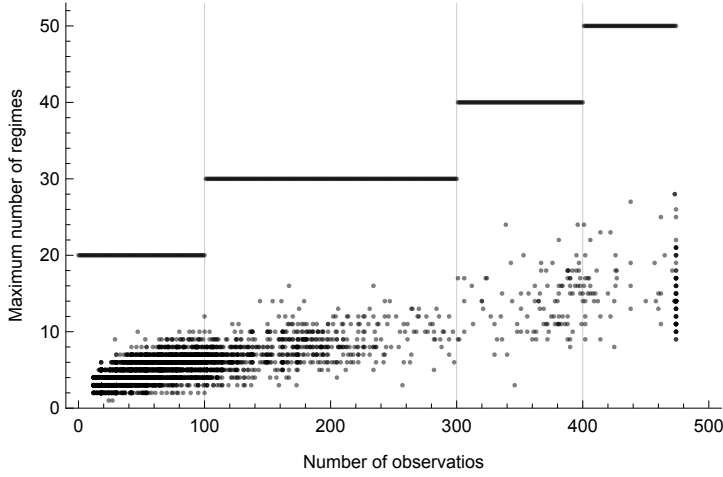


Figure 1: The four horizontal lines are the pre-set thresholds for the maximum number of in- and out-of-sample regimes,  $M_i$ , plotted against the length of the fund’s history,  $\mathcal{T}_i$ . Each point is the largest MCMC draw of  $s_{i,T_i}$  under the nonparametric priors.

global hyperparameter for all the fund-regimes, whereas the nonparametric model can draw a different hyperparameter value for each fund-regime.

For both models we start collecting the posterior draws after discarding the first 30,000 MCMC draws. We run the sampler for the nonparametric priors first and use its last draw to initialize the unknowns in the parametric hierarchical prior. Noticeable trending was observed in the initial draws of the nonparametric approach but by the end of the burnin period the sampler had settled down and was generating well behaved draws of the unknowns. We then iterate each sampler for 60,000 more sweeps keeping every sixtieth draw to construct a random sample of 1,000 posterior draws.<sup>18</sup>

### 5.3 Number of regimes

In both the nonparametric and parametric case the maximum number of fund-regimes,  $M_i$ , is set so as not to restrict the number of in-sample regimes. We use four different values, 20, 30, 40 and 50 fund-regimes where the funds with the shortest histories have  $M_i = 20$  and those with the longest have  $M_i = 50$  (this is visually shown in Figure 1). By discriminating on the number of returns we are able to speed up the sampler.

To put to rest any doubts that our posterior results have been affected by the preset value of the  $M_i$ , in Figure 1 we plot  $M_i$  and the largest draw of  $s_{i,T_i}$  under the nonparametric

<sup>18</sup>Because the functional form of the conditional posterior densities have known analytical formulas, 1,000 draws is large enough to accurately represent these posterior densities after marginalizing out the unknowns over the 1,000 draws.

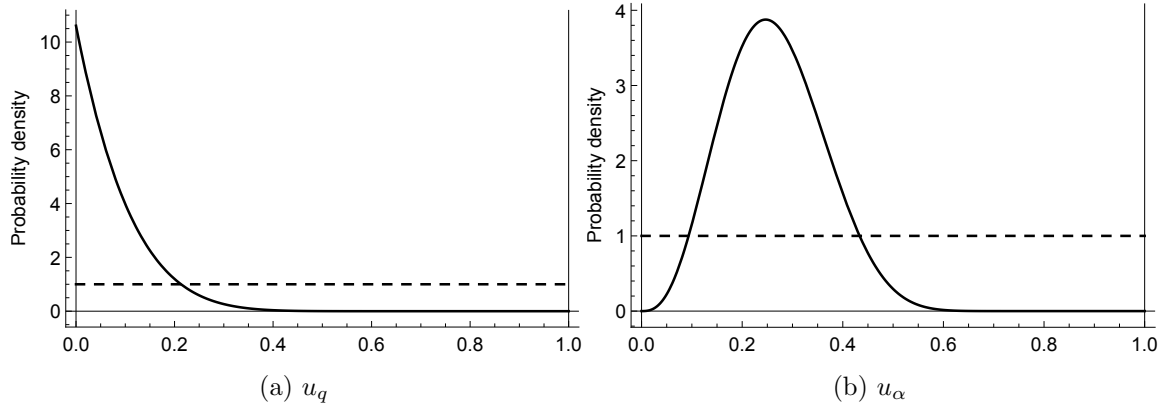


Figure 2: Posterior distribution of the transformed concentration parameter  $u_q$  in panel (a) and  $u_\alpha$  in panel (b).

prior against  $\mathcal{T}_i$  (the technical online appendix contains the same figure for the parametric prior). In the figures the largest draw of  $s_{i,T_i}$  never gets close to  $M_i$ . It is interesting to point out that the maximum draw of  $s_{i,T_i}$  range from two to almost thirty regimes. This suggests Bollen & Whaley (2009) single-change-point model may be too restrictive.

Under the nonparametric priors there are an average of 12,390 different fund-regimes over all the funds, which produces an average of 2.4 regimes per fund. The average fund-regime then lasts around three years. Under the parametric priors the mean number of total fund-regimes drops to 11,914. This drop is statistically different since the posterior densities of the number of fund-regimes from the two approaches do not overlap.<sup>19</sup> Hence, we conclude that the parametric priors lead to fewer fund-regimes than the nonparametric priors.

## 5.4 Model comparison

Because the nonparametric hierarchical prior nests the parametric when the concentration parameter is zero, the Savage-Dickey (SD) density ratio of Dickey (1971) can be used to quickly compute the Bayes factor in favor of the parametric prior ( $\eta_\bullet = 0$ ). In terms of the transformed concentration parameter,  $u_\bullet = \eta_\bullet / (1 + \eta_\bullet)$ , the SD density ratio is

$$\text{BF}(u_\bullet = 0) \equiv \frac{\pi(u_\bullet = 0|\mathcal{Y})}{\text{Uniform}(u_\bullet|0, 1)},$$

where  $\pi(u_\bullet|\mathcal{Y})$  is the posterior distribution of  $u_\bullet$ .

<sup>19</sup>The posterior densities of the number of fund-regimes can be found in the online Appendix.

In Figure 2 we plot in panel (a) the Rao-Blackwellized posterior distribution of the transformed concentration parameters  $u_q$  for the transition probabilities,  $q_i$ , and in panel (b), the posterior of  $u_\alpha$ . The uniform prior for the  $u_\bullet$ s are indicated in the plots by the horizontal dashed line at one. Since the prior equals one over its range of zero to one, visually a SD ratio in favor of the parametric prior is the vertical height of the posterior at zero. In panel (a) we see that the odds of the data supporting the parametric hierarchical prior of  $q_i$  are slightly better than ten-to-one over the nonparametric prior. On the other hand, in panel (b) the odds in favor of the parametric distribution of alpha is essentially zero.

Panel (b) of Figure 2 shows there is strong empirical evidence in favor of the distribution of the alphas having more than one cluster. The posterior is effectively zero at  $u_\alpha = 0$ , meaning that a parametric hierarchical prior of the alphas is a poor modeling choice. The density is also effectively zero at  $u_\alpha = 1$  so there is essentially zero empirical evidence for each fund-regime having its own unique cluster. Given the near zero probabilities for  $u_\alpha$  equaling zero or one, assuming the hyperprior,  $G_\alpha$ , is known apriori would be a bad modeling assumption.

## 5.5 Posterior predictive distributions

In Figures 3, 4, and 5, we plot the nonparametric posterior predictive densities,  $\pi_q(q|\mathcal{Y})$ , for the transition probabilities,  $\pi_\alpha(\alpha|\mathcal{Y})$ , for the alphas, and,  $\pi_\sigma(\sigma^2|\mathcal{Y})$ , for the variances, respectively. Each figure includes the corresponding predictive density under the parametric hierarchical priors. Eq. (25) is used to compute the density in Figure 5 and similar Rao-Blackwell formulas are used to compute the other posterior predictive densities.

The nonparametric predictive distribution of the transition probability plotted in Figure 3 has a posterior median of one cluster; ie.,  $\mathcal{K}_q = 1$ . Given the height of the density in panel (a) of Figure 2 at the origin, finding there to be only one cluster is not a surprise. Having only one cluster complete sharing of information under both the nonparametric and parametric prior of  $q$  occurs.

Both posterior predictive densities of  $q$  strongly favor a multiple-change-point process for the parameters of the four-factor risk model. The probability of  $q$  being less than 0.006 is essentially zero under both posterior densities of Figure 3. Therefore, skill clearly changes over time and a mutual fund performance model whose alpha is assumed to be constant would amount to a very restrictive prior.

The nonparametric density  $\pi_\alpha(\alpha|\mathcal{Y})$  plotted in Figure 4 has a posterior median of four clusters,  $\mathcal{K}_\alpha = 4$ , with three clear modes at approximately  $-1.75$ ,  $1.5$ , and  $6$ . Since our

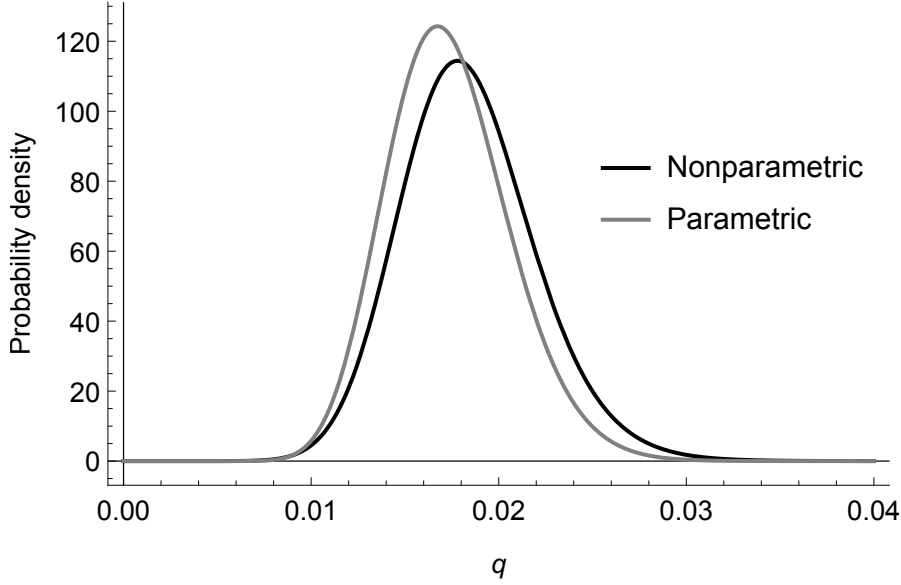


Figure 3: The parametric and nonparametric posterior predictive density of the transition probability,  $q$ ,  $\pi_q(q|\mathcal{Y})$ .

goal differs from Malsiner-Walli et al. (2016), and, instead, focuses on flexibly modeling the distribution of the fund-regime alphas, we are unable to make economic statements about these modes.<sup>20</sup> Nonetheless, these modes suggest the average performance levels of three distinct groups, unskilled, break-even, and skilled fund-regimes, respectively.

The break-even mode of 1.5 has economic support in Berk & Green (2004) who find the average actively managed mutual fund charges a fee of 1.5 percent. The secondary modes are also consistent with Berk & Green (2004) and their decreasing return to scale model of fund performance. Under their model an alpha greater than the break-even level leads to investors increasing their investment in the skilled fund. This causes the assets under management of the fund to increase and, because of the decreasing returns to scale, causes its alpha to move to a lower fund-regime. The opposite occurs for an unskilled fund.

Alphas belonging to the same cluster shrink towards the cluster average. In Figure 4 the greatest shrinkage is found in the tight variance around the  $\alpha_{ims}$  belonging to the hypothesized break-even cluster. In contrast, the parametric predictive density shows complete sharing and global shrinkage. Alphas from the extraordinary fund-regimes end up being

<sup>20</sup>This is because of identification issues that arise from the label switching problem detailed by Geweke (2007) and Frühwirth-Schnatter (2006).

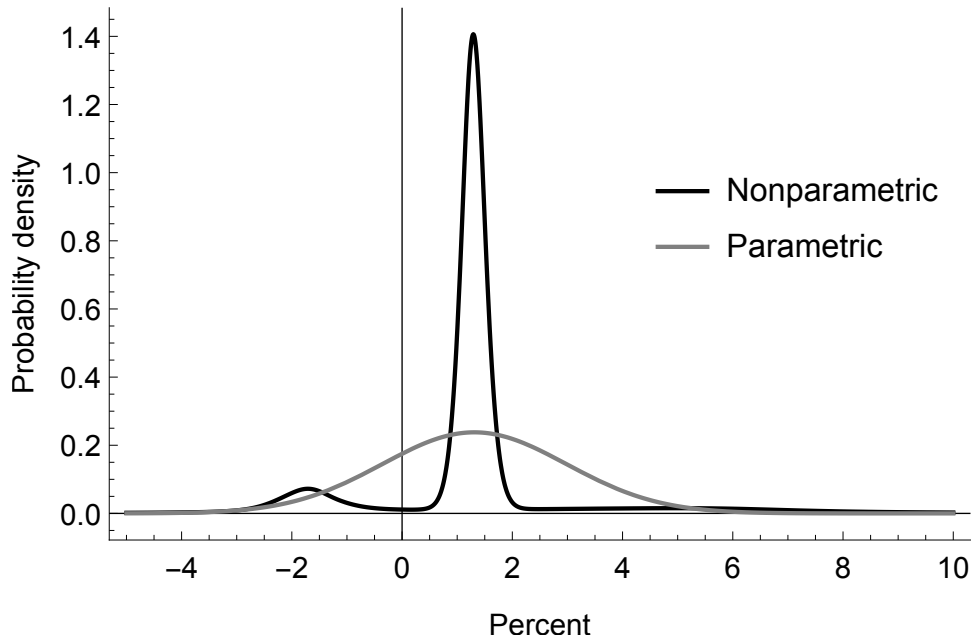


Figure 4: Parametric and nonparametric posterior predictive density of alpha for a future fund-regime,  $\pi_\alpha(\alpha|\mathcal{Y})$ .

contaminated by the information from the average fund-regime skill level of the entire population and end looking like the break-even regime.

In Figure 5 we plot the nonparametric posterior predictive density of the regime variances. This density has two modes – a primary mode near a variance of one-hundred and a secondary mode near one.<sup>21</sup> Except for the mode near one, the parametric posterior predictive density looks identical to the nonparametric density.

Our working hypothesis for the mode near one is that it belongs to a group of index funds that track the market. By tracking market returns these funds are less risky and have on average a lower variance. In contrast non-market index funds take greater risks in order to outperform the market.

## 5.6 Posterior fund performance

Each mutual fund in our panel has its own history of posterior smoothed dynamic densities for  $\alpha_{i,t}$ . Of these we have selected the highest performing fund with the longest tenure to illustrate the smoothed densities of the alphas. The fund is the Fidelity Magellan Fund which opened for business in 1963. Magellan had one of the best known fund manager in

<sup>21</sup>Please refer to the online Appendix for a graph that more clearly shows the mode near zero.

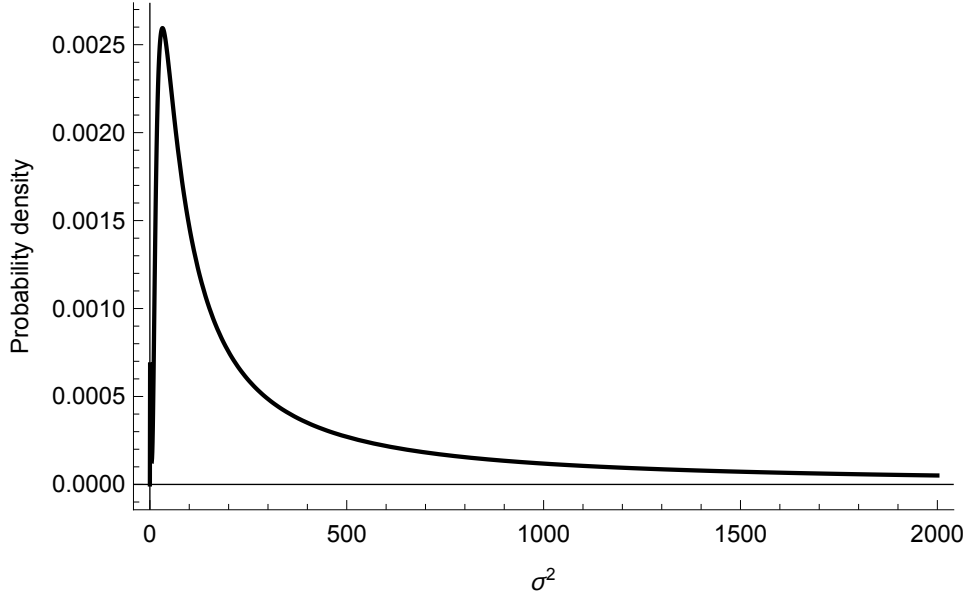


Figure 5: Posterior nonparametric predictive density of a future fund-regimes variance,  $\pi_\sigma(\sigma^2|\mathcal{Y})$ .

Peter Lynch, who managed the fund for thirteen years from May 1977 to May 1990. Besides Lynch the fund had been managed by five other fund managers over our sample period.

Under the nonparametric priors Magellan experiences on average sixteen regimes over its entire time period, whereas the parametric priors uncovers an average of seventeen regimes. Magellan’s probability of switching regimes under both types of hierarchical priors is nearly identical. Under the nonparametric priors the transition probability is 0.022, whereas it is 0.021 for the parametric priors. Both probabilities are only slightly larger than the 0.018 posterior transition probability for the population. Thus, regardless of the hierarchical prior, the duration of a Magellan regime is only slightly longer than the typical fund-regime.

In Figure 6 we plot the posterior median and highest probability intervals (HPD) for the 90, 75, and 50 percentile of Magellan’s posterior distribution  $\pi_\mu(\alpha_{FM,t}|\mathcal{Y})$  from June 1963 to June 2001 (please refer to the online technical appendix for the calculation of these distributions). In the top panel are the HPDs from the nonparametric priors and in the bottom panel are those for the parametric priors. Each fund manager’s tenure is denoted by the five vertical lines with the longest tenure being that of Lynch.

Comparing the top and bottom panels of Figure 6 we see how the global shrinkage caused by the parametric (bottom) approach dampens the exceptional regimes relative to

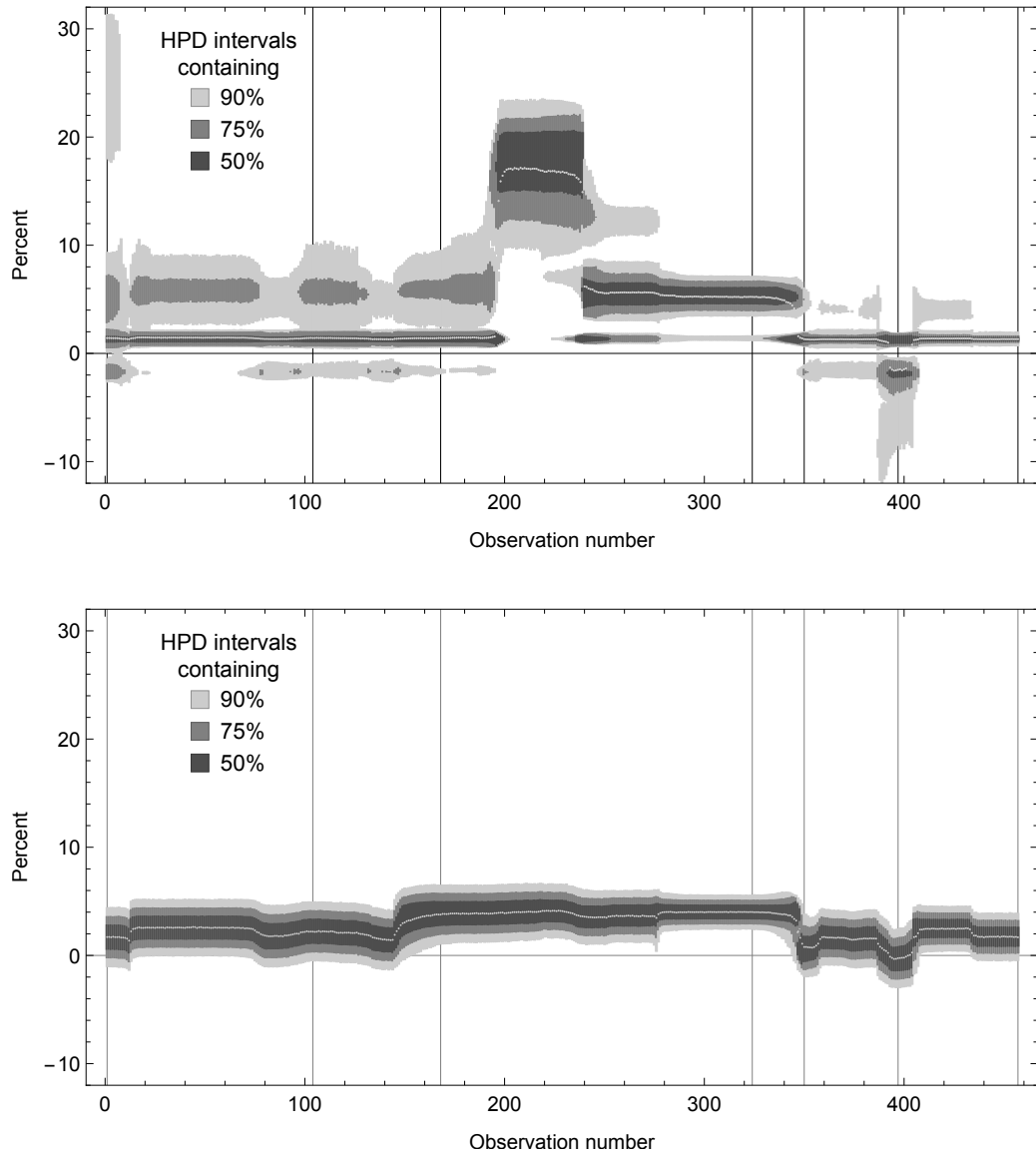


Figure 6: The highest probability density intervals (HPD) and posterior medians (white dots) under the nonparametric prior (top graph) and parametric prior (bottom graph) from the posterior densities  $\pi(\alpha_{FM,t}|\mathcal{Y})$ ,  $t = 06/1963, \dots, 06/2011$  for Fidelity Magellan. The vertical lines at 104 (01/1972), 168 (05/1977), 324 (05/1990), 350 (07/1992), 397 (06/1996), and 457 (06/2001) denote a change in the fund manager. Lynch was the manager from 168 to 324.

the local shrinkage of the nonparametric priors (top). Although the parametric posteriors during the Lynch era (observations 168 to 324) had higher medians and HPD intervals than during the rest of Magellan’s history, the intervals shifted up before Lynch. Contrast this with the dramatic increase in the nonparametric posterior’s medians and HPDs after Lynch took over. Under the parametric prior any exceptional performance attributable to Lynch is muted and occurred before he became manager.

Except for the instance where the alphas for Magellan were already transitioning to a new regime, changing managers fails to result in any change in fund performance under either the parametric or nonparametric priors. Both panels in Figure 6 suggest Magellan continued to deliver the same alpha before and after a change in the manager.

The nonparametric posterior medians in Figure 6 begin transitioning later and take fewer months to reach their new regime than do the parametric posterior medians. Local shrinkage under the nonparametric approach is the likely reason. Instead of shrinking towards the global average, local shrinkage pulls the medians towards the local average. This explains the quick transition from Lynch’s early ordinary performance levels to his extraordinary returns of eighteen to nineteen percent a year.

Figure 7 plots each of the posterior densities  $\pi_\mu(\alpha_{FM,t}|\mathcal{Y})$  used to compute the HPDs in the top panel of Figure 6. Magellan’s five fund manager’s tenures are denoted by different gray scales along with axis ticks. Many of the densities resemble the nonparametric posterior predictive density plotted in Figure 4 with their three modes.<sup>22</sup> This resemblance occurs because the posterior predictive distribution is essentially the prior for Magellan once the prior predictive has been updated with the information contained in the return data of the 5,135 other funds.

There are also densities in Figure 7 that are visibly different from the population distribution of the alphas. These exceptional densities have fewer than three modes and are centered at larger values of alpha. For instance, during the Lynch period the posterior is unimodal at 18%. These episodes are periods of truly exceptional skill and performance. Lynch’s investment strategy during this period eliminated the negative and break-even modes. This exceptional performance lasts for nearly three years before the break-even mode reappears and the primary mode retreats to a lower but still exceptional alpha of six percent. This extraordinary shape, relative to the population distribution of skill, is kept for the remainder of Lynch’s time as manager and through most of the tenure of the next manager.

---

<sup>22</sup>Although not shown here this is also the case for Magellan’s parametric posterior densities,  $\pi_\mu(\alpha_{FM,T_{FM}+1}|\mathcal{Y})$ .

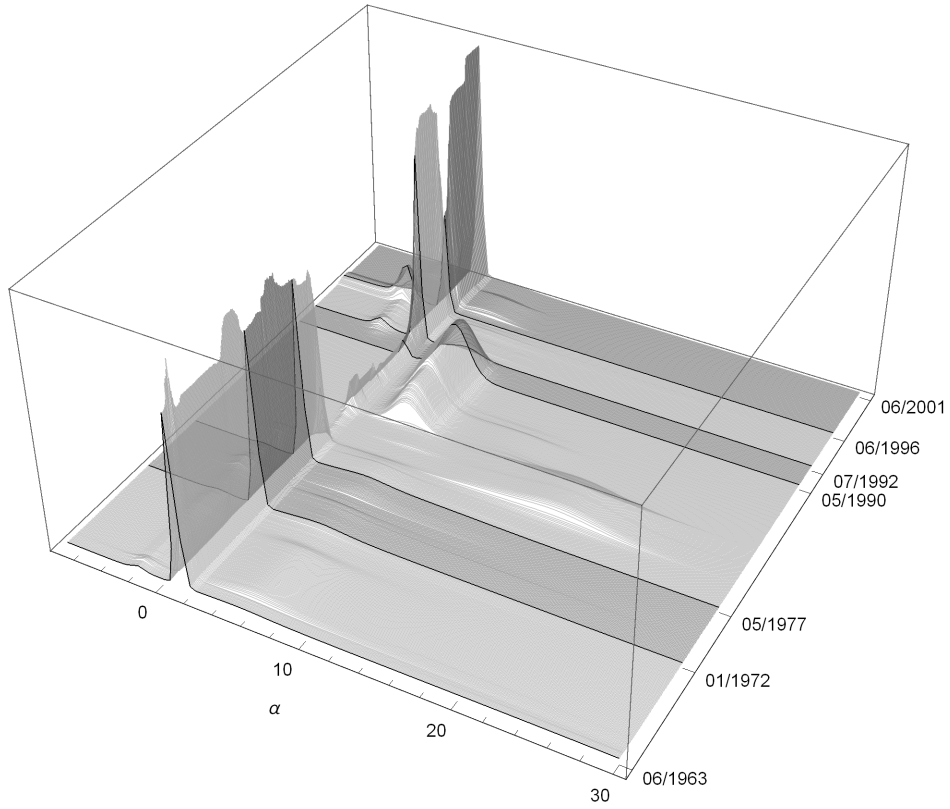


Figure 7: Posterior densities for the alpha of the Fidelity Magellan Fund,  $\pi(\alpha_{FM,t}|\mathcal{Y})$ , where  $t = 06/1963, \dots, 06/2001$ . A change in the shade of gray denotes the end to, and the start of, one of Magellan's six fund managers. These dates are denoted by tics on the x-axis (Lynch's tenure began in 05/1977 and ended in 05/1990).

## 6 Conclusion

In this paper we extended the change point model to a panel of multiple-change-point processes where the parameters are modeled with nonparametric hierarchical priors. Under nonparametric hierarchical priors our multiple-change-point panel model shares information from across the regimes and panel individuals through the hierarchical prior’s hyperparameters and their unknown hyperprior distribution to produce more robust parameter estimates. Our nonparametric hierarchical priors cluster together the individuals and their regimes into groups that have similar behavior and locally shrink the unknown parameters towards the average of the group. This partial sharing and local shrinkage allows extreme regimes and individuals to be identified.

We apply our multiple-change-point panel model and nonparametric hierarchical priors to a longitudinal data set of mutual fund returns to investigate the skill level of mutual funds and how likely they are to maintain this skill into the future. There is overwhelming empirical evidence supporting change-point behavior in the fund-regime parameters. Hence, assuming that mutual fund skill is constant over the history of a fund is a very bad assumption.

Under the nonparametric hierarchical priors we find on average four clusters and three modes for the population distribution of skill; a primary mode at 1.5 percent a year, another mode at  $-1.75$ , and a very diffuse mode at around 6 percent. Compared to the parametric prior’s single mode at 1.5 percent there is strong empirical evidence against using the parametric hierarchical prior to estimate skill. Having three modes keeps the estimate of skilled and unskilled fund’s performance from being masked by shrinking to the average performance level of 1.5 percent.

Our findings for the population carries over to the performance of one of the longest lived skilled funds in the panel, Fidelity Magellan. We find Magellan to have extended periods of time where its performance is average, highly skilled, and also low skilled. Under the parametric prior these periods of extraordinary performance were not that extraordinary. Periods of exceptional skill were hidden by the parametric prior’s restrictive single mode assumption for the population distribution of skill.

In future research we aim to design a method capable of finding the funds and their regimes that have similar posterior densities and use them to investigate the structural nature of skill. These and many other interesting questions can be addressed using our nonparametric hierarchical prior with the multiple-change-point panel model.

## References

- Bassetti, F., Casarin, R. & Leisen, F. (2014), ‘Beta-product dependent Pitman–Yor processes for Bayesian inference’, *Journal of Econometrics* **180**(1), 49 – 72.
- Bauwens, L., Koop, G., Korobilis, D. & Rombouts, J. V. (2015), ‘The contribution of structural break models to forecasting macroeconomic series’, *Journal of Applied Econometrics* **30**(4), 596–620.
- Berk, J. B. & Green, R. C. (2004), ‘Mutual fund flows and performance in rational markets’, *Journal of Political Economy* **112**, 1269–1295.
- Billio, M., Casarin, R., Ravazzolo, F. & Van Dijk, H. K. (2016), ‘Interconnections between Eurozone and US booms and busts using a Bayesian panel Markov-switching VAR model’, *Journal of Applied Econometrics* **31**(7), 1352–1370.
- Bollen, N. P. B. & Busse, J. A. (2004), ‘Short-term persistence in mutual fund performance’, *The Review of Financial Studies* **18**(2), 569–597.
- Bollen, N. P. B. & Whaley, R. E. (2009), ‘Hedge fund risk dynamics: Implications for performance appraisal’, *The Journal of Finance* **64**(2), 985–1035.
- Busse, J. A. & Irvine, P. J. (2006), ‘Bayesian alphas and mutual fund persistence’, *The Journal of Finance* **61**(5), 2251–2288.
- Carhart, M. M. (1997), ‘On persistence in mutual fund performance’, *The Journal of Finance* **52**(1), pp. 57–82.
- Chan, J. C. & Koop, G. (2014), ‘Modelling breaks and clusters in the steady states of macroeconomic variables’, *Computational Statistical and Data Analysis* **76**, 186–193.
- Chib, S. (1996), ‘Calculating posterior distributions and modal estimates in Markov mixture models’, *Journal of Econometrics* **75**, 79–97.
- Chib, S. (1998), ‘Estimation and comparison of multiple change point models’, *Journal of Econometrics* **86**, 221–241.
- Chib, S. & Hamilton, B. (2002), ‘Semiparametric Bayes analysis of longitudinal data treatment models’, *Journal of Econometrics* **110**, 67–89.
- Chib, S. & Tiwari, R. (1988), ‘Bayes prediction density and regression estimation: A semi parametric approach’, *Empirical Economics* **13**, 209–222.

- Dickey, J. D. (1971), ‘The weighted likelihood ratio, linear hypotheses on normal location parameters’, *Annals of Mathematical Statistics* **42**, 204–223.
- Dubey, S. D. (1970), ‘Compound gamma, beta and f distributions’, *Metrika* **16**(1), 27–31.
- Dufays, A. (2015), ‘Infinite-state Markov-switching for dynamic volatility’, *Journal of Financial Econometrics* **14**(2), 418.
- Elton, E. J. & Gruber, M. J. (2013), Mutual funds, *in* M. H. George M. Constantinides & R. M. Stulz, eds, ‘Handbook of the Economics of Finance’, Vol. 2, Part B, Elsevier, pp. 1011 – 1061.
- Elton, E. J., Gruber, M. J. & Blake, C. R. (1996), ‘The persistence of risk-adjusted mutual fund performance’, *Journal of Business* **69**(2), 133–157.
- Escobar, M. D. & West, M. (1995), ‘Bayesian density estimation and inference using mixtures’, *Journal of the American Statistical Association* **90**(430), 577–588.
- Fama, E. F. & French, K. R. (1993), ‘Common risk factors in the returns on stocks and bonds’, *Journal of Financial Economics* **33**(1), 3 – 56.
- Ferguson, T. (1973), ‘A Bayesian analysis of some nonparametric problems’, *The Annals of Statistics* **1**(2), 209–230.
- Fisher, M. (2017), Nonparametric density estimation using a mixture of order-statistics distributions, Technical report, Federal Reserve Bank of Atlanta.  
**URL:** <http://www.markfisher.net/mefisher/papers/multi-resolution.pdf>
- Frühwirth-Schnatter, S. (2006), *Finite Mixture and Markov Switching Models*, Springer.
- Frühwirth-Schnatter, S. & Kaufmann, S. (2006), ‘How do changes in monetary policy affect bank lending? An analysis of Austrian bank data’, *Journal of Applied Econometrics* **21**(3), 275–305.
- Frühwirth-Schnatter, S. & Kaufmann, S. (2008), ‘Model-based clustering of multiple time series’, *Journal of Business & Economic Statistics* **26**(1), 78–89.
- Gelman, A., Carlin, J. B., Stern, H. S., Dunson, D. B., Vehtari, A. & Rubin, D. B. (2013), *Bayesian Data Analysis*, Chapman and Hall/CRC.
- Gerlach, R., Carter, C. & Kohn, R. (2000), ‘Efficient Bayesian inference for dynamic mixture models’, *Journal of the American Statistical Association* **95**(451), 819–828.

- Geweke, J. (2007), ‘Interpretation and inference in mixture models: Simple MCMC works’, *Computational Statistics & Data Analysis* **51**(7), 3529 – 3550.
- Geweke, J. & Jiang, Y. (2011), ‘Inference and prediction in a multiple-structural-break model’, *Journal of Econometrics* **163**, 172–185.
- Giordani, P. & Kohn, R. (2008), ‘Efficient Bayesian inference for multiple change-point and mixture innovation models’, *Journal of Business & Economic Statistics* **26**(1), 66–77.
- Goetzmann, W. N. & Ibbotson, R. G. (1994), ‘Do winners repeat? Patterns in mutual fund return behavior’, *Journal of Portfolio Management* **20**(2), 9–18.
- Grinblatt, M. & Sheridan, T. (1992), ‘The persistence of mutual fund performance’, *The Journal of Finance* **47**(5), 1977–1984.
- Hendricks, D., Patel, J. & Zeckhauser, R. (1993), ‘Hot hands in mutual funds: Short-run persistence of relative performance, 1974 – 1988’, *The Journal of Finance* **48**(1), 93–130.
- Hirano, K. (2002), ‘Semiparametric Bayesian inference in autoregressive panel data models’, *Econometrica* **70**, 781–799.
- Hjort, N. L., Holmes, C., Mueller, P. & Walter, S. G., eds (2010), *Bayesian Nonparametrics*, Cambridge University Press.
- Jensen, M. C. (1968), ‘The performance of mutual funds in the period 1945-1964’, *The Journal of Finance* **23**(2), pp. 389–416.
- Jensen, M. J. & Maheu, J. M. (2010), ‘Bayesian semiparametric stochastic volatility modeling’, *Journal of Econometrics* **157**(2), 306 – 316.
- Jin, X. & Maheu, J. M. (2016), ‘Bayesian semiparametric modeling of realized covariance matrices’, *Journal of Econometrics* **192**(1), 19 – 39.
- Jochmann, M. (2015), ‘Modeling U.S. inflation dynamics: A Bayesian nonparametric approach’, *Econometric Reviews* **34**(5), 537–558.
- Jones, C. S. & Shanken, J. (2005), ‘Mutual fund performance with learning across funds’, *Journal of Financial Economics* **78**, 507–552.
- Koijen, R. S. (2014), ‘The cross-section of managerial ability, incentives, and risk preferences’, *The Journal of Finance* **69**(3), 1051–1098.

- Koop, G. & Potter, S. M. (2007), ‘Estimation and forecasting in models with multiple breaks’, *Review of Economic Studies* **74**, 763–789.
- Lijoi, A., Prünster, I. & Walker, S. G. (2005), ‘On consistency of nonparametric normal mixtures for Bayesian density estimation’, *Journal of the American Statistical Association* **100**(472), 1292–1296.
- MacEachern, S. N. & Müller, P. (1998), ‘Estimating mixture of Dirichlet process models’, *Journal of Computational and Graphical Statistics* **7**(2), 223–238.
- Maheu, J. M. & Song, Y. (2014), ‘A new structural break model, with an application to Canadian inflation forecasting’, *International Journal of Forecasting* **30**(1), 144 – 160.
- Malsiner-Walli, G., Frühwirth-Schnatter, S. & Grün, B. (2016), ‘Model-based clustering based on sparse finite Gaussian mixtures’, *Statistics and Computing* **26**(1), 303–324.
- McCulloch, R. E. & Tsay, R. (1993), ‘Bayesian inference and prediction for mean and variance shifts in autoregressive time series’, *Journal of the American Statistical Association* **88**, 968–978.
- Neal, R. (2000), ‘Markov chain sampling methods for Dirichlet process mixture models’, *Journal of Computational and Graphical Statistics* **9**, 249–265.
- Pesaran, M. H., Pettenuzzo, D. & Timmermann, A. (2006), ‘Forecasting time series subject to multiple structural breaks’, *Review of Economic Studies* **73**(4), 1057 – 1084.
- Petrone, S. & Wasserman, L. (2002), ‘Consistency of Bernstein polynomial posteriors’, *Journal of the Royal Statistical Society: Series B (Statistical Methodology)* **64**(1), 79–100.
- Song, Y. (2014), ‘Modelling regime switching and structural breaks with an infinite hidden Markov model’, *Journal of Applied Econometrics* **29**(5), 825–842.
- Teh, W. Y. (2010), Dirichlet process, in C. Sammut & G. I. Webb, eds, ‘Encyclopedia of Machine Learning’, Springer US, Boston, MA, chapter Dirichlet Process, pp. 280–287.
- West, M., Muller, P. & Escobar, M. (1994), Hierarchical priors and mixture models with applications in regression and density estimation, in P. R. Freeman & A. F. Smith, eds, ‘Aspects of Uncertainty’, John Wiley.

# Online Appendix

## A The model

In this section we describe a factor model of returns for *fund-regimes*. Within a single fund-regime, the model is quite standard. The novelty lies in two places: the determination of the regimes within a single fund and the connections across the fund-regimes via hyperparameters in the prior distribution.

### Fund returns and regimes

Let  $J$  denote the number of mutual funds and let  $i \in \{1, \dots, J\}$  index the funds. For fund  $i$  there are  $\mathcal{T}_i$  observations that occur at times

$$t \in \mathcal{T}_i = \{\tau_i, \tau_i + 1, \dots, T_i\}, \quad (\text{A.1})$$

where  $1 \leq \tau_i < T_i \leq T_{\max}$ . Note  $\mathcal{T}_i = |\mathcal{T}_i|$ . Let  $y_{i,t}$  denote the  $t$ -th observation for the  $i$ -th fund. Let

$$Y_{i,t} := (y_{i\tau_i}, \dots, y_{it}), \quad (\text{A.2})$$

so that  $Y_{i,T_i}$  denotes the full set of returns for fund  $i$ .

Let  $M_i$  denote the maximum number of regimes for fund  $i$  and let  $m \in \{1, \dots, M_i\}$  index the regime number. Let  $s_{i,t}$  denote the regime number for fund  $i$  at observation  $t$ . Note  $s_{i,T_i} \leq \min[\mathcal{T}_i, M_i]$ . Let

$$S_{i,t} := (s_{i\tau_i}, \dots, s_{it}), \quad (\text{A.3})$$

so that  $S_{i,T_i}$  denotes the full set of regime indicators for fund  $i$ . Let  $\mathcal{S} := \{S_{i,T_i}\}_{i=1}^J$  denote the complete set of states for all funds. We make the following independence assumption regarding the distribution of the observations given the states:

$$p(\{Y_{i,T_i}\}_{i=1}^J | \mathcal{S}) = \prod_{i=1}^J p(Y_{i,T_i} | S_{i,T_i}) = \prod_{i=1}^J \prod_{t=\tau_i}^{T_i} p(y_{i,t} | s_{i,t}). \quad (\text{A.4})$$

### Assigned and unassigned regimes

The index set of potential regimes is

$$\mathcal{J}_{\mathcal{M}} := \{(i, m) : i \in \{1, \dots, J\} \wedge m \in \{1, \dots, M_i\}\}. \quad (\text{A.5})$$

The number of potential regimes is  $\mathcal{N}_{\mathcal{M}} = \sum_{i=1}^J M_i$ . Conditioning on the states, the index set of *assigned* regimes (i.e., those regimes associated with observations) is

$$\mathcal{J}_{\mathcal{A}} := \{(i, m) \in \mathcal{J}_{\mathcal{M}} : m \leq s_{i, T_i}\}, \quad (\text{A.6})$$

which has  $\mathcal{N}_{\mathcal{A}} = \sum_{i=1}^J s_{i, T_i}$  elements. The index set of unassigned regimes is

$$\mathcal{J}_{\mathcal{U}} := \{(i, m) \in \mathcal{J}_{\mathcal{M}} : m > s_{i, T_i}\} = \mathcal{J}_{\mathcal{M}} \setminus \mathcal{J}_{\mathcal{A}}, \quad (\text{A.7})$$

which has  $\mathcal{N}_{\mathcal{U}} = \mathcal{N}_{\mathcal{M}} - \mathcal{N}_{\mathcal{A}}$  elements.<sup>23</sup>

## A single fund-regime

The model is connected to the data via the fund-regimes. Let

$$\mathcal{Y}_{i, m} := \{y_{i, t} \in Y_{i, T_i} : s_{i, t} = m\} \quad (\text{A.8})$$

denote the set of observations for regime  $m$  in fund  $i$  and let  $\mathcal{T}_{i, m}$  denote the number of observations in  $\mathcal{Y}_{i, m}$ . The fund-regime constitutes the unit of analysis. For unassigned regimes, where  $m > s_{i, T_i}$ , we have  $\mathcal{Y}_{i, m} = \emptyset$  and  $\mathcal{T}_{i, m} = 0$ .

We restrict ourselves to factor models of the following sort. Let  $N$  denote the number factors in the model of returns. For  $y_{i, t} \in \mathcal{Y}_{i, m}$  given  $m \leq s_{i, T_i}$ ,

$$y_{i, t} = \alpha_{i, m} + \sum_{j=1}^N \beta_{i, m}^j F_{it}^j + \varepsilon_{i, t}, \quad (\text{A.9a})$$

where

$$\varepsilon_{i, t} \stackrel{\text{iid}}{\sim} \mathbf{N}(0, \sigma_{i, m}^2). \quad (\text{A.9b})$$

For more compact notation, let

$$\phi_{im} := (\phi_{im}^0, \phi_{im}^1, \dots, \phi_{im}^N) = (\alpha_{i, m}, \beta_{i, m}^1, \dots, \beta_{i, m}^N) \quad (\text{A.10a})$$

$$\theta_{im} := (\phi_{im}, \sigma_{i, m}^2). \quad (\text{A.10b})$$

In addition, let  $X_{it} = (1, F_{it}^1, \dots, F_{it}^N)$ , where  $F_{it} = (F_{it}^1, \dots, F_{it}^N)$  is the vector of market-wide factors properly-aligned to fund  $i$ . With this notation, we can express (A.9) as

$$p(y_{i, t} | s_{i, t} = m) = p(y_{i, t} | \theta_{im}) = \mathbf{N}(y_{i, t} | X_{it}^\top \phi_{im}, \sigma_{i, m}^2). \quad (\text{A.11})$$

---

<sup>23</sup>One way to eliminate structural breaks is to set  $M_i = 1$  for all  $i$ , in which case  $\mathcal{J}_{\mathcal{M}} = \mathcal{J}_{\mathcal{A}}$ ,  $\mathcal{J}_{\mathcal{U}} = \emptyset$ , and  $\mathcal{N}_{\mathcal{M}} = J$ . Another way is to set the probability of a regime change to zero for all funds. See below.

Consequently, the likelihood for  $\theta_{im}$  can be expressed as

$$p(\mathcal{Y}_{im}|\theta_{im}) = \mathbf{N}(\mathcal{Y}_{im}|\mathcal{X}_{im}\phi_{im}, \sigma_{i,m}^2 I_{\mathcal{T}_{im}}), \quad (\text{A.12})$$

where

$$\mathcal{X}_{im} := \{X_{it} : s_{i,t} = m\}. \quad (\text{A.13})$$

We complete the model of a fund-regime by providing a prior for the fund-regime parameters. The prior for  $\theta_{im}$  (conditional on hyperparameters) is<sup>24</sup>

$$p(\theta_{im}|a_{im}, h_{\phi_{im}}^2, h_{\sigma_{im}}^2) = \mathbf{N}(\phi_{im}|a_{im}, B_{im}) \text{Inv-Gamma}(\sigma_{im}^2|\nu_{\sigma}/2, h_{\sigma_{im}}^2 \nu_{\sigma}/2), \quad (\text{A.14})$$

where  $B_{im} = \text{diag}(h_{\phi_{im}}^2)$  is a  $(N+1) \times (N+1)$  diagonal matrix and

$$a_{im} := (a_{im}^0, a_{im}^1, \dots, a_{im}^N) \quad (\text{A.15a})$$

$$h_{\phi_{im}}^2 := (h_{\phi_{im}}^{2,0}, h_{\phi_{im}}^{2,1}, \dots, h_{\phi_{im}}^{2,N}). \quad (\text{A.15b})$$

### Reexpressing the hyperparameters

Before we proceed, it is convenient to express the hyperparameters in the prior for  $\theta_{im}$  [see (A.15)] as follows. Let

$$\psi_{im}^j := (a_{im}^j, h_{\phi_{im}}^{2,j}) \quad \text{and} \quad \psi_{im}^{\sigma} := h_{\sigma_{im}}^2. \quad (\text{A.16})$$

Using this notation, we can reexpress the prior for  $\theta_{im}$  [see (A.14)] as

$$\phi_{im}^j|\psi_{im}^j \sim \mathbf{N}(a_{im}^j, h_{\phi_{im}}^{2,j}) \quad \text{for } j = 0, \dots, N \quad (\text{A.17a})$$

$$\sigma_{im}^2|\psi_{im}^{\sigma} \sim \text{Inv-Gamma}(\nu_{\sigma}/2, h_{\sigma_{im}}^2 \nu_{\sigma}/2). \quad (\text{A.17b})$$

This formulation will prove useful in computing the posterior distributions for the hyperparameters.

---

<sup>24</sup>For reference, note

$$\begin{aligned} \text{Inv-Gamma}(\sigma^2|\nu/2, s^2 \nu/2) &= \frac{e^{-\frac{\nu s^2}{2\sigma^2}} \left(\frac{\nu s^2}{2\sigma^2}\right)^{\nu/2}}{s^2 \Gamma(\frac{\nu}{2})} \\ \text{Gamma}(s^2|a, 1/b) &= \frac{e^{-b s^2} (b s^2)^{a-1}}{b^{-1} \Gamma(a)}. \end{aligned}$$

## Prior for the regimes

We complete the model of a fund by providing a prior for the regime changes within a fund.

Let  $q_i$  denote the probability of a change in regime (i.e., state) for the  $i$ -th fund. The conditional prior for the regimes within a fund is

$$p(S_{iT_i}|q_i) = p(s_{i,\tau_i}) \prod_{t=\tau_i}^{T_i-1} p(s_{i,t+1}|s_{it}, q_i), \quad (\text{A.18})$$

where  $p(s_{i,\tau_i} = 1) = 1$  and

$$p(s_{i,t+1} = \ell | s_{i,t} = m, q_i) = \begin{cases} 1 - q_i & \ell = m < M_i \\ q_i & \ell = m + 1 \\ 1 & \ell = m = M_i \\ 0 & \text{otherwise} \end{cases}. \quad (\text{A.19})$$

The prior for  $q_i$  is given by

$$p(q_i | \psi_i^q) = \text{Beta}(q_i | j_i, k_i - j_i + 1), \quad (\text{A.20})$$

where

$$\psi_i^q := (j_i, k_i) \quad (\text{A.21})$$

denotes the hyperparameter in the prior for  $q_i$ . We assume conditional independence across funds as follows:

$$p(\{(S_{iT_i}, q_i)\}_{i=1}^J | \{\psi_i^q\}_{i=1}^J) = \prod_{i=1}^J p(S_{iT_i} | q_i) p(q_i | \psi_i^q). \quad (\text{A.22})$$

## Priors for the hyperparameters

A central feature of the model is the sharing of hyperparameters across fund-regimes. The mechanism for sharing involves grouping hyperparameters into clusters and assigning a common value to all hyperparameters within a given cluster — the value of the *cluster parameter*. In other words, we identify a fund-regime parameter with the parameter of the cluster to which it has been assigned. We adopt the Chinese Restaurant Process (CRP) representation of the Dirichlet process prior distribution for the clustering. The cluster parameters themselves have a prior distribution called the *base distribution* (which the fund-regime hyperparameters inherit).

We first provide a skeleton for the hyperprior that can be applied to the various cases.

## Hyperprior skeleton

We begin with the classifications. The CRP can be characterized as follows:

$$z_{1:K}|\eta \sim \text{CRP}(\eta), \quad (\text{A.23})$$

where  $z_{1:K}$  denotes a collection of the  $K$  classifications variables,  $z_{\kappa}$ ,  $\kappa = 1, \dots, K$ , and  $\eta$  denotes the concentration parameter of the Dirichlet process  $DP(\eta, G_0)$ . In order to provide an explicit representation, let  $\mathcal{K}_{\kappa} = \max[z_{1:\kappa}]$ , the number of clusters in  $z_{1:\kappa}$ , and let  $n_c^{\kappa}$  denote the multiplicity of  $c$  in  $z_{1:\kappa}$  so that  $\sum_{c=1}^{\mathcal{K}_{\kappa}} n_c^{\kappa} = \mathcal{K}_{\kappa}$ . Then (A.23) can be expressed as

$$p(z_{1:\mathcal{K}}|\eta) = \prod_{\kappa=1}^{K-1} p(z_{\kappa+1}|z_{1:\kappa}, \eta), \quad (\text{A.24})$$

where  $z_1 = 1$  and

$$p(z_{\kappa+1} = c|z_{1:\kappa}, \eta) = \begin{cases} \frac{n_c^{\kappa}}{\mathcal{K}_{\kappa} + \eta} & c \in \{1, \dots, \mathcal{K}_{\kappa}\} \\ \frac{\eta}{\mathcal{K}_{\kappa} + \eta} & c = \mathcal{K}_{\kappa} + 1 \end{cases}. \quad (\text{A.25})$$

Having determined the classifications, a cluster parameter is drawn from the base distribution  $G_0$  for each of the  $\mathcal{K}_K$  clusters:

$$\psi_c^* \stackrel{\text{iid}}{\sim} G_0. \quad (\text{A.26})$$

Finally, combining the classifications with the cluster parameters, a hyperparameter is set equal to the cluster parameter for the cluster to which it (the hyperparameter) belongs:

$$\psi_{\kappa} = \psi_{z_{\kappa}}^*. \quad (\text{A.27})$$

## Details for clustering

We put a prior on all of the hyperparameters associated with the full set of potential factor-model coefficients; namely,

$$\{\psi_{im}^j\}_{\mathcal{M}} := \{\psi_{im}^j : (i, m) \in \mathcal{J}_{\mathcal{M}}\} \quad (\text{A.28a})$$

$$\{\psi_{im}^{\sigma}\}_{\mathcal{M}} := \{\psi_{im}^{\sigma} : (i, m) \in \mathcal{J}_{\mathcal{M}}\}. \quad (\text{A.28b})$$

Let  $z_{im}^j$  denote the cluster to which  $\psi_{im}^j$  is assigned, so that  $z_{im}^j = c$  means  $\psi_{im}^j = \psi_c^{j*}$ . Similarly, let  $z_{im}^{\sigma}$  denote the cluster assignment for  $\psi_{im}^{\sigma}$ . The prior for the cluster assignments can be expressed as

$$\{z_{im}^j\}_{\mathcal{M}}|\eta_j \sim \text{CRP}(\eta_j) \quad (\text{A.29a})$$

$$\{z_{im}^{\sigma}\}_{\mathcal{M}}|\eta_{\sigma} \sim \text{CRP}(\eta_{\sigma}). \quad (\text{A.29b})$$

where  $\eta_j$  and  $\eta_\sigma$  are the DP prior's *concentration parameters*.

Let  $z_i^q$  denote the cluster assignment for  $\psi_i^q$  and let  $\{z_i^q\}_{1:J}$  denote the set of assignments for all funds. The prior for the classifications can be expressed as

$$\{z_i^q\}_{1:J} | \eta_q \sim \text{CRP}(\eta_q), \quad (\text{A.30})$$

where  $\eta_q$  is the concentration parameter.

The prior for the concentration parameters is given by

$$\log(\eta_\ell) \sim \text{Logistic}(1, 1), \quad (\text{A.31})$$

for  $\ell \in \{j, \sigma, q\}$ .<sup>25</sup>

### Base distributions

We now turn to the base distributions for the cluster parameters. Let us denote the cluster parameters by

$$\psi_c^{j*} = (a_c^{j*}, h_{\phi c}^{2j*}), \quad \psi_c^{\sigma*} = h_{\sigma c}^{2*}, \quad \text{and} \quad \psi_c^{q*} = (j_c^*, k_c^*), \quad (\text{A.32})$$

where  $c$  is an index for the cluster number. (Each of the cluster parameters has its own index  $c$ .) The base distributions for  $\psi_c^{j*}$  and  $\psi_c^{\sigma*}$  are given by

$$p(\psi_c^{j*}) = p(a_c^{j*}, h_{\phi c}^{2j*}) = \text{N}(a_c^{j*} | a_0^j, h_{\phi c}^{2j*} / \kappa_0^j) \text{Inv-Gamma}(h_{\phi c}^{2j*} | \nu_0^j / 2, h_0^{2j} \nu_0^j / 2) \quad (\text{A.33a})$$

and

$$p(\psi_c^{\sigma*}) = p(h_{\sigma c}^{2*}) = \text{Gamma}(h_{\sigma c}^{2*} | c_0, 1/b_0). \quad (\text{A.33b})$$

The base distribution for  $\psi_c^{q*}$  can be expressed as  $p(\psi_c^{q*}) = p(j_c^*, k_c^*) = p(j_c^* | k_c^*) p(k_c^*)$ , where<sup>26</sup>

$$j_c^* | k_c^* \sim \text{Uniform}(1, \dots, k_c^*) \quad (\text{A.34a})$$

$$k_c^* - 1 \sim \text{Geometric}(\xi_0). \quad (\text{A.34b})$$

---

<sup>25</sup>For reference, note  $\text{Log-Logistic}(x|a, b) = ab^{-a}x^{a-1}/(1 + (x/b)^a)^2$ .

<sup>26</sup>Let  $\xi_0 = 1/200$ .

## B Sampling from the posterior

We now present our approach to sampling from the posterior distribution. We divide the parameters into blocks and sample each block conditional on the parameters in the other blocks (Gibbs sampling approach). For each block we adopt an appropriate strategy for sampling.

There are two types of blocks: within-fund blocks and across-fund blocks. The within-fund blocks are conditionally independent across funds and consequently can be handled in parallel. The across-fund blocks involve drawing the hyperparameters (via the classifications and the cluster parameters) and the concentration parameters.

### Fund-regime parameters $\{\theta_{im}\}$

Given the likelihood (A.12) and the conditionally conjugate prior (A.14), the two conditional posteriors are given by

$$p(\phi_{im}|\mathcal{Y}_{im}, \sigma_{im}^2, a_{im}, B_{im}) = \mathbf{N}(\phi_{im}|\bar{a}_{im}, \bar{B}_{im}) \quad (\text{B.1a})$$

where

$$\bar{B}_{im} = \left( \sigma_{im}^{-2} \mathcal{X}_{im}^\top \mathcal{X}_{im} + B_{im}^{-1} \right)^{-1} \quad (\text{B.1b})$$

$$\bar{a}_{im} = \bar{B}_{im} \left( \sigma_{im}^{-2} \mathcal{X}_{im}^\top \mathcal{Y}_{im} + B_{im}^{-1} a_{im} \right) \quad (\text{B.1c})$$

and

$$p(\sigma_{im}^2|\mathcal{Y}_{im}, \phi_{im}, h_{\sigma im}^2) = \text{Inv-Gamma}(\sigma_{im}^2|\bar{\nu}_{im}/2, \bar{h}_{\sigma im}^2 \bar{\nu}_{im}/2) \quad (\text{B.2a})$$

where

$$\bar{\nu}_{im} = \nu_\sigma + \mathcal{T}_{im} \quad (\text{B.2b})$$

$$\bar{h}_{\sigma im}^2 = \frac{h_{\sigma im}^2 \nu_\sigma + (\mathcal{Y}_{im} - \mathcal{X}_{im} \phi_{im})^\top (\mathcal{Y}_{im} - \mathcal{X}_{im} \phi_{im})}{\bar{\nu}_{im}}. \quad (\text{B.2c})$$

For unassigned regimes (those regimes for which  $\mathcal{Y}_{im} = \emptyset$ ), the posterior reduces to the prior (conditional on the hyperparameters). [See (A.14).]

### Regimes $\{S_{iT_i}\}$

Factor the joint posterior distribution of the states as follows:

$$p(S_{iT_i}|Y_{i,T_i}, \Theta_i, q_i) = p(s_{i,T_i}|Y_{i,T_i}, \Theta_i, q_i) p(s_{i,T_i-1}|S_i^{T_i}, Y_{i,T_i}, \Theta_i, q_i) \cdots \\ p(s_{i,t}|S_i^{t+1}, Y_{i,T_i}, \Theta_i, q_i) \cdots p(s_{i,1}|S_i^2, Y_{i,T_i}, \Theta_i, q_i), \quad (\text{B.3})$$

where

$$\Theta_i = (\theta_{i1}, \dots, \theta_{iM_i}) \quad (\text{B.4a})$$

$$S_i^t = (s_{i,t}, \dots, s_{i,T_i}). \quad (\text{B.4b})$$

Draws of the states will be made using (B.3) starting with  $s_{i,T_i}$  and working backwards.

The conditional distributions in (B.3) can be computed as follows:

$$p(s_{i,t}|S_i^{t+1}, Y_{i,T_i}, \Theta_i, q_i) \propto p(s_{i,t}|Y_{i,t}, \Theta_i, q_i) p(s_{i,t+1}|s_{i,t}, q_i), \quad (\text{B.5})$$

where

$$p(s_{i,t} = m|Y_{i,t}, \Theta_i, q_i) = \frac{p(s_{i,t} = m|Y_{i,t-1}, \Theta_i, q_i) p(y_{i,t}|\theta_{im})}{\sum_{m'=1}^{M_i} p(s_{i,t} = m'|Y_{i,t-1}, \Theta_i, q_i) p(y_{i,t}|\theta_{i,m'})} \quad (\text{B.6})$$

and where

$$\begin{aligned} p(s_{i,t} = m|Y_{i,t-1}, \Theta_i, q_i) \\ = \sum_{m'=m-1}^m p(s_{i,t} = m|s_{i,t-1} = m', q_i) p(s_{i,t-1} = m'|Y_{i,t-1}, \Theta_i, q_i) \end{aligned} \quad (\text{B.7})$$

for  $t \geq 2$  and  $p(s_{i,1} = 1|Y_{i0}, \Theta_i, q_i) = 1$ .

Let us streamline the notation and bring out the recursive structure. We define three matrices. First, let  $Q_i$  denote the  $M_i \times M_i$  matrix where [see (A.19)]

$$(Q_i)_{m\ell} = p(s_{i,t+1} = \ell|s_{i,t} = m, q_i). \quad (\text{B.8})$$

Second, let  $L^i$  denote the  $T_i \times M_i$  matrix where

$$L_{tm}^i = p(y_{i,t}|\theta_{im}) = \mathbf{N}(y_{i,t}|X_{it}^\top \phi_{im}, \sigma_{i,m}^2), \quad (\text{B.9})$$

and let  $L_t^i$  denote the  $t$ th row of  $L^i$ . Third, let  $G^i$  denote the  $T_i \times M_i$  matrix where<sup>27</sup>

$$G_{tm}^i = p(s_{i,t} = m|Y_{i,t-1}, \Theta_i, q_i) p(y_{i,t}|\theta_{im}) \quad (\text{B.10})$$

and similarly let  $G_t^i$  denote the  $t$ th row of  $G^i$ .

With this notation, the forward recursion is given by

$$G_t^i \propto (G_{t-1}^i Q_i) \circ L_t^i, \quad (\text{B.11})$$

---

<sup>27</sup>Note  $G_{tm}^i = p(s_{i,t} = m, y_{i,t}|Y_{i,t-1}, \Theta_i, q_i)$ .

where “ $\circ$ ” denotes the Hadamard (component-by-component) product. The recursion is initialized with  $G_1^i = (L_{11}^i, 0, \dots, 0)$ . The probabilities of the states at  $t = T_i$  are given by

$$p(s_{i,T_i} = m | Y_{i,T_i}, \Theta_i, q_i) \propto G_{T_i m}^i, \quad (\text{B.12})$$

while for  $t < T_i$  they are

$$p(s_{i,t} = m | s_{i,t+1} = \ell, Y_{i,T_i}, \Theta_i, q_i) \propto (G_t^i \circ Q_i^\ell)_m, \quad (\text{B.13})$$

where  $Q_i^\ell$  denotes the  $\ell$ th column of  $Q_i$ .<sup>28</sup>

The backward sampler begins with (B.12) and continues with (B.13) iteratively until  $s_{i,1}$  is drawn.

### Regime-change probabilities $\{q_i\}$

Given the states and the classifications, we can draw  $q_i$  as follows. The likelihood for  $q_i$  is

$$\text{Binomial}(s_{i,T_i} - 1 | \mathcal{T}_i - 1, q_i), \quad (\text{B.14})$$

since  $s_{i,T_i} - 1$  is the number of change-points in  $T_i - 1$  independent Bernoulli trials where  $q_i$  is the probability of a change-point for a single Bernoulli trial. Therefore the conditional posterior distribution is

$$p(q_i | s_{i,T_i}, \psi_i^q) = \text{Beta}(q_i | j_i + s_{i,T_i} - 1, k_i - j_i + 1 + (\mathcal{T}_i - s_{i,T_i})). \quad (\text{B.15})$$

### Hyperparameters

We now turn to drawing the hyperparameters. Recall there are three sets of hyperparameters:  $\{\psi_{im}^j\}_{\mathcal{M}}$ ,  $j = 0, 1, \dots, N$ ,  $\{\psi_{im}^\sigma\}_{\mathcal{M}}$ , and  $\{\psi_i^q\}_{1:J}$ , where  $\psi_{im}^j = (a_{im}^j, h_{\phi im}^{2j})$ ,  $\psi_{im}^\sigma = h_{\sigma im}^2$ , and  $\psi_i^q = (j_i, k_i)$ .

Regarding  $\psi_{im}^j$  and  $\psi_{im}^\sigma$ , it is necessary to distinguish between those hyperparameters associated with regimes that have been assigned (conditional on the states  $\mathcal{S}$ ) and those that have not. In particular, let  $\{\psi_{im}^j\}_{\mathcal{A}}$  and  $\{\psi_{im}^\sigma\}_{\mathcal{A}}$  denote the sets of these hyperparameters for regimes that have been assigned and let  $\{\psi_{im}^j\}_{\mathcal{U}}$  and  $\{\psi_{im}^\sigma\}_{\mathcal{U}}$  denote the sets of these hyperparameters for regimes that are unassigned.

For each of  $\{\psi_{im}^j\}_{\mathcal{A}}$ ,  $\{\psi_{im}^\sigma\}_{\mathcal{A}}$ , and  $\{\psi_i^q\}_{1:J}$ , we adopt the approach to sampling based on Algorithm 2 in Neal (2000), although the details differ across the three sets of hyperparameters. For the remaining two sets of hyperparameters,  $\{\psi_{im}^j\}_{\mathcal{U}}$  and  $\{\psi_{im}^\sigma\}_{\mathcal{U}}$ , we sample from the distributions conditioned on the respective assigned hyperparameters  $\{\psi_{im}^j\}_{\mathcal{A}}$  and  $\{\psi_{im}^\sigma\}_{\mathcal{A}}$ .

---

<sup>28</sup>If  $s_{i,t} = 1$  then  $s_{i,t-1} = 1$ , while if  $s_{i,t} = t$ , then  $s_{i,t-1} = t - 1$ . In either of these cases no more computation is required.

### Drawing $\{\psi_{im}^j\}_{\mathcal{A}}$

We must establish some additional notation. Define the set of indices of specific regimes associated with cluster  $c$ ,

$$\mathcal{I}_c^j := \{(i, m) \in \mathcal{I}_{\mathcal{A}} : z_{im}^j = c\}, \quad (\text{B.16})$$

and the set of parameters associated with cluster  $c$ ,

$$\{\phi_{im}^j\}_c := \{\phi_{im}^j | (i, m) \in \mathcal{I}_c^j\}. \quad (\text{B.17})$$

Let  $n_c^j$  denote the number of elements in  $\{\phi_{im}^j\}_c$ . Let  $\{z_{im}^j\}_{\mathcal{A}}$  denote the set of classifications for assigned regimes. Assume  $\{z_{im}^j\}_{\mathcal{A}}$  is normalized so that the set of unique values equals  $\{1, 2, \dots, \mathcal{K}_j\}$ , where  $\mathcal{K}_j$  is the number of distinct classifications. Note  $\sum_{c=1}^{\mathcal{K}_j} n_c^j = \mathcal{N}_{\mathcal{A}}$ .

Note that  $\{\phi_{im}^j\}_c$  plays the role of the “observations” with regard to the “parameter”  $\psi_c^{j*}$  in the conditional posterior. In this role, the likelihood for a single observation is given by the prior (A.17a). Thus the posterior for  $\psi_c^{j*}$  given  $\{\phi_{im}^j\}_c$  is:

$$p(\psi_c^{j*} | \{\phi_{im}^j\}_c) = \mathbf{N}(a_c^{j*} | \bar{a}_j, h_{\phi_c}^{2j*} / \bar{\kappa}_j) \text{Inv-Gamma}(h_{\phi_c}^{2j*} | \bar{\nu}_j / 2, \bar{h}_j^2 \bar{\nu}_j / 2), \quad (\text{B.18})$$

where

$$\bar{\kappa}_j = \kappa_0^j + n_c^j \quad (\text{B.19a})$$

$$\bar{\nu}_j = \nu_0^j + n_c^j \quad (\text{B.19b})$$

$$\bar{a}_j = \left( \frac{\kappa_0^j}{\bar{\kappa}_j} \right) a_0^j + \left( \frac{n_c^j}{\bar{\kappa}_j} \right) \bar{\phi}_c^j \quad (\text{B.19c})$$

$$\bar{h}_j^2 = \left( \frac{\nu_0^j}{\bar{\nu}_j} \right) h_0^{j2} + \left( \frac{n_c^j}{\bar{\nu}_j} \right) \hat{\sigma}_{\phi_c^j}^2 + \left( \frac{\kappa_0^j}{\bar{\kappa}_j} \right) \left( \frac{n_c^j}{\bar{\nu}_j} \right) (\bar{\phi}_c^j - a_0^j)^2, \quad (\text{B.19d})$$

where the “observations” are summarized by

$$\bar{\phi}_c^j := \frac{1}{n_c^j} \sum_{\iota \in \mathcal{I}_c^j} \phi_{\iota}^j \quad \text{and} \quad \hat{\sigma}_{\phi_c^j}^2 := \frac{1}{n_c^j} \sum_{\iota \in \mathcal{I}_c^j} (\phi_{\iota}^j - \bar{\phi}_c^j)^2. \quad (\text{B.20})$$

### Classification

It remains to describe how the elements of  $\{z_{im}^j\}_{\mathcal{A}}$  (the classifications for the assigned regimes) are drawn. Let  $\{z_{im}^j\}_{\mathcal{A}}^{-\iota}$  denote the (possibly renormalized) vector of classifications after having removed case  $\iota$  for some  $\iota \in \mathcal{I}_{\mathcal{S}}$ . Renormalization will occur if — before removal — the cluster associated with observation  $\iota$  is a singleton (i.e.,  $n_{z_{\iota}^j}^j = 1$ ), in which case  $\psi_{z_{\iota}^j}^{j*}$

will be discarded, the remaining clusters will be relabeled, and the corresponding classifications will be adjusted. Let  $\mathcal{K}_j^{-\iota} = \max[\{z_{im}^j\}_{\mathcal{A}}^{-\iota}]$  and let  $(n_c^j)^{-\iota}$  denote the multiplicity of class  $c$  in  $\{z_{im}^j\}_{\mathcal{A}}^{-\iota}$ .

The full conditional probability of  $z_{\iota}^j$  given the “observations” is

$$p(z_{\iota}^j | \{z_{im}^j\}_{\mathcal{A}}^{-\iota}, \phi_{\iota}^j, \eta_j) \propto \begin{cases} \frac{(n_c^j)^{-\iota}}{N_{\mathcal{A}} - 1 + \eta_j} p(\phi_{\iota}^j | z_{\iota}^j = c) & c \in \{1, \dots, \mathcal{K}_j^{-\iota}\} \\ \frac{\eta_j}{N_{\mathcal{A}} - 1 + \eta_j} p(\phi_{\iota}^j) & c = \mathcal{K}_j^{-\iota} + 1 \end{cases}, \quad (\text{B.21})$$

where

$$p(\phi_{im}^j | z_{im}^j = c) = \mathbf{N}(\phi_{im}^j | a_c^{j*}, h_{\phi_c}^{2j*}), \quad (\text{B.22})$$

and

$$\begin{aligned} p(\phi_{im}^j) &= \int \mathbf{N}(\phi_{im}^j | a_c^{j*}, h_{\phi_c}^{2j*}) \mathbf{N}(a_c^{j*} | a_0^j, h_{\phi_c}^{2j*} / \kappa_0^j) \text{Inv-Gamma}(h_{\phi_c}^{2j*} | \nu_0^j / 2, h_0^{j2} \nu_0^j / 2) d\psi_c^{j*} \\ &= t_{\nu_0^j}(\phi_{im}^j | a_0^j, (1 + 1/\kappa_0^j) h_0^{j2}). \end{aligned} \quad (\text{B.23})$$

If a new cluster is chosen, then populate it with a draw from the posterior  $\psi_c^{j*} | \phi_{im}^j$  [see (B.18)].

### Drawing $\{\psi_{im}^{\sigma}\}_{\mathcal{A}}$

Let

$$\mathcal{I}_c^{\sigma} := \{(i, m) \in \mathcal{I}_{\mathcal{A}} : z_{im}^{\sigma} = c\}, \quad (\text{B.24})$$

and let  $\{\sigma_{im}^2\}_c := \{\sigma_{im}^2 | (i, m) \in \mathcal{I}_c^{\sigma}\}$ , where the number observations in  $\{\sigma_{im}^2\}_c$  equals  $n_c^{\sigma}$ . Let  $\mathcal{K}_{\sigma}$  denote the number of distinct classifications in the collection  $\{z_{im}^{\sigma}\}$ .

Again,  $\{\sigma_{im}^2\}_c$  play the role of the observations where the likelihood for a single observation is given by (A.17b). The posterior for  $h_{\sigma c}^{2*}$  given  $\{\sigma_{im}^2\}_c$  is

$$p(h_{\sigma c}^{2*} | \{\sigma_{im}^2\}_c) = \text{Gamma}(h_{\sigma c}^{2*} | \bar{c}_c, 1/\bar{b}_c), \quad (\text{B.25})$$

where

$$\bar{c}_c = c_0 + \frac{\nu_{\sigma}}{2} n_c^{\sigma} \quad (\text{B.26a})$$

$$\bar{b}_c = b_c + \frac{\nu_{\sigma}}{2} \sum_{\iota \in \mathcal{I}_c^{\sigma}} \frac{1}{\sigma_{\iota}^2}. \quad (\text{B.26b})$$

## Classification

The probabilities for classification involve the following likelihoods: For existing clusters,

$$p(\sigma_{im}^2 | z_{im}^\sigma = c) = \text{Inv-Gamma}(\sigma_{im}^2 | \nu_\sigma/2, h_{\sigma c}^{2*} \nu_\sigma/2), \quad (\text{B.27})$$

while for a new cluster we use the predictive distribution:

$$\begin{aligned} p(\sigma_{im}^2) &= \int \text{Inv-Gamma}(\sigma_{im}^2 | \nu_\sigma/2, h_{\sigma c}^{2*} \nu_\sigma/2) \text{Gamma}(h_{\sigma c}^{2*} | c_0, 1/b_0) dh_{\sigma c}^{2*} \\ &= \frac{\left(\frac{\sigma_{im}^2}{w} + 1\right)^{-(c_0 + \frac{\nu_\sigma}{2})} \left(\frac{\sigma_{im}^2}{w}\right)^{c_0-1}}{w B(c_0, \frac{\nu_\sigma}{2})} \quad \text{where } w = \nu_\sigma/(2b_0), \end{aligned} \quad (\text{B.28})$$

which is the PDF of the three-parameter generalized Beta-Prime distribution.<sup>29</sup> If a new cluster is chosen, then populate it with a draw from the posterior  $h_{\sigma c}^{2*} | \sigma_{im}^2$  [see (B.25)].

## Unassigned regimes

Regarding  $\{\psi_{im}^j\}_U$  and  $\{\psi_{im}^\sigma\}_U$ , we draw the classifications from the CRP conditioned on the cluster counts of the assigned hyperparameters:

$$\{z_{im}^j\}_U | \{z_{im}^j\}_A, \eta_j \sim \text{CRP}(\{n_c^j\}_{c=1}^{\mathcal{K}_j}, \eta_j) \quad (\text{B.29a})$$

$$\{z_{im}^\sigma\}_U | \{z_{im}^\sigma\}_A, \eta_\sigma \sim \text{CRP}(\{n_c^\sigma\}_{c=1}^{\mathcal{K}_\sigma}, \eta_\sigma), \quad (\text{B.29b})$$

where  $\text{CRP}(\{n_c\}_{c=1}^{\mathcal{K}}, \eta)$  indicates conditioning on the existing classifications. The unassigned hyperparameters that are classified with already existing clusters are set equal to the cluster parameters from those clusters. For each new cluster drawn via (B.29), a cluster parameter is drawn from the associated base distribution [see (A.33)].

## Drawing $\{\psi_i^q\}_{1:J}$

We adopt the scheme described in Fisher (2017) (see for omitted details), which uses Neal's Algorithm 2.

Let

$$\mathcal{I}_c^q := \{i \in \{1, \dots, J\} : z_i^q = c\}, \quad (\text{B.30})$$

and let  $\{q_i\}_c := \{q_i | i \in \mathcal{I}_c^q\}$ , where the number observations in  $\{q_i\}_c$  equals  $n_c^q$ . Let  $\mathcal{K}_q$  denote the number of distinct classifications in the collection  $\{z_i^q\}$ .

<sup>29</sup>Since  $\text{Inv-Gamma}(\sigma_{im}^2 | \nu_\sigma/2, h_{\sigma c}^{2*} \nu_\sigma/2) \equiv \text{Gamma}(h_{\sigma c}^{2*} | \nu_\sigma/2 + 1, 2\sigma_{im}^2/\nu_\sigma)$ , (B.28) is also the PDF of a ‘‘compound gamma’’ distribution.

Fisher (2017) presents two methods for drawing  $\psi_c^{q*}$ . The first method assumes  $n_c^q = 1$ . In this case, the conditional posterior is given by

$$p(j_c^*, k_c^* | q_i) = \frac{p(q_i | j_c^*, k_c^*) p(j_c^* | k_c^*) p(k_c^*)}{p(q_i)} = p(j_c^* | k_c^*, q_i) p(k_c^*), \quad (\text{B.31})$$

where

$$p(j_c^* | k_c^*, q_i) = \binom{k_c^* - 1}{j_c^* - 1} q_i^{j_c^* - 1} (1 - q_i)^{k_c^* - j_c^*} = \text{Binomial}(j_c^* - 1 | k_c^* - 1, q_i). \quad (\text{B.32})$$

Note that  $k_c^*$  is not identified and is drawn from its prior distribution and then  $(j_c^* - 1) \sim \text{Binomial}(k_c^* - 1, q_i)$  with the proviso that if  $k_c^* = 1$  then  $j_c^* = 1$ .

Now assume  $n_c^q \geq 2$ . The posterior for  $(j_c^*, k_c^*)$  is given by

$$\begin{aligned} p(\psi_c^{q*} | \{q_i\}, \{z_i^q\}, \eta) &= p(\psi_c^{q*} | \{q_i\}_c) \\ &= p(j_c^*, k_c^*) \times \frac{(\prod_{i \in \mathcal{I}_c^q} q_i)^{j_c^* - 1} (\prod_{i \in \mathcal{I}_c^q} (1 - q_i))^{k_c^* - j_c^*}}{B(j_c^*, k_c^* - j_c^* + 1)^{n_c^q}}. \end{aligned} \quad (\text{B.33})$$

One can adopt a Metropolis-Hastings scheme for the  $(r + 1)$  draw given the  $r$ th draw with the following proposal:

$$k_c^{*'} - 1 \sim \text{Poisson}(k_c^{*(r)}) \quad (\text{B.34})$$

$$j_c^{*'} - 1 \sim \text{Binomial}(k_c^{*'} - 1, \bar{q}_c^{(r)}), \quad (\text{B.35})$$

where  $\bar{q}_c = \frac{1}{n_c^q} \sum_{i \in \mathcal{I}_c^q} q_i$  is the average of  $\{q_i\}_c$ . Let

$$\widehat{p}((\psi_c^{q*})' | \psi_c^{q*}, \bar{q}_c) := \text{Poisson}(k_c^{*'} - 1 | k_c^*) \text{Binomial}(j_c^{*'} | k_c^{*'} - 1, \bar{q}_c). \quad (\text{B.36})$$

Then

$$\psi_c^{q*(r+1)} = \begin{cases} (\psi_c^{q*})' & \mathcal{M}_c^{(r)} \geq u^{(r+1)}, \\ \psi_c^{q*(r)} & \text{otherwise,} \end{cases} \quad (\text{B.37})$$

where  $u^{(r+1)} \sim \text{Uniform}(0, 1)$  and

$$\mathcal{M}_c^{(r)} = \frac{p((\psi_c^{q*})' | \{q_i\}_c^{(r)})}{p((\psi_c^{q*})^{(r)} | \{q_i\}_c^{(r)})} \times \frac{\widehat{p}((\psi_c^{q*})^{(r)} | (\psi_c^{q*})', \bar{q}_c^{(r)})}{\widehat{p}((\psi_c^{q*})' | (\psi_c^{q*})^{(r)}, \bar{q}_c^{(r)})}. \quad (\text{B.38})$$

## Classification

The full conditional probability of  $z_i$  given the “observations” is

$$p(z_i | \eta_q, \{z_i^q\}^{-i}, \{\psi_c^*\}^{-i}, q_i) \propto \begin{cases} \frac{(n_c^q)^{-i}}{j-1+\eta_q} \text{Beta}(q_i | j_c^*, k_c^* - j_c^* + 1) & c \in \{1, \dots, \mathcal{K}_q^{-i}\} \\ \frac{\eta_q}{j-1+\eta_q} \text{Uniform}(q_i | 0, 1) & c = \mathcal{K}_q^{-i} + 1 \end{cases}. \quad (\text{B.39})$$

The uniform predictive for a new cluster follows from the adding-up property of Bernstein polynomials in conjunction with the uniform conditional prior for  $j_i^*|k_i^*$ .

If a new cluster is selected, then the new cluster parameter is drawn from the posterior using the scheme described above for  $n_i^q = 1$ .

### Concentration parameters

For the purposes of this section, let  $\eta$  stand for any of the concentration parameters ( $\eta_j$ ,  $\eta_\tau$ , or  $\eta_q$ ). In addition, let  $\{z_i\}$  denote the corresponding set of classifications, let  $n_c$  denote the number of elements in cluster  $c$ , and let  $\mathcal{K}$  denote the number of clusters. Let  $B = \sum_{c=1}^{\mathcal{K}} n_c$ . (Note,  $B = \mathcal{N}_s$  for  $\eta_j$  and  $\eta_\sigma$  while  $B = J$  for  $\eta_q$ .) It can be shown that the likelihood for the concentration parameter is

$$p(\{z_i\}|\eta) = \frac{\eta^{\mathcal{K}} \prod_{c=1}^{\mathcal{K}} (n_c - 1)!}{(\eta)_B} \propto \frac{\eta^{\mathcal{K}}}{(\eta)_B}, \quad (\text{B.40})$$

where  $(\eta)_B := \prod_{j=1}^B (j - 1 + \eta) = \Gamma(B + \eta)/\Gamma(\eta)$ .

It is convenient to change variables to  $w = \log(\eta)$ . Recall  $w \sim \text{Logistic}(1, 1)$ . We can express the product of the likelihood and the prior in terms of  $w$ :

$$p(\{z_i\}|w)p(w) = \frac{e^{(\mathcal{K}+1)w} \Gamma(e^w)}{(1 + e^w)^2 \Gamma(B + e^w)}. \quad (\text{B.41})$$

Draws of  $w$  can be made via a Metropolis scheme using a symmetric proposal distribution such as  $w' \sim \text{Logistic}(w^{(r)}, h)$ , where  $w^{(r)} = \log(\eta^{(r)})$ . The next element in the chain is given by

$$\eta^{(r+1)} = \begin{cases} e^{w'} & \frac{p(\{z_i\}|w')p(w')}{p(\{z_i\}|w^{(r)})p(w^{(r)})} \geq u^{(r+1)}, \\ \eta^{(r)} & \text{otherwise,} \end{cases} \quad (\text{B.42})$$

where  $u^{(r+1)} \sim \text{Uniform}(0, 1)$ .

## C Posterior distributions

In this section we discuss the posterior distributions for the population (generic) and specific mutual fund (specific) cases.

### Population distribution of the factor coefficients

Let  $\mathbb{Y} := \{Y_{iT_i}\}_{i=1}^J$ . For each factor coefficient there is a posterior predictive distribution for the generic case of a  $g$ th unobserved mutual fund:

$$p(\phi_g^j|\mathbb{Y}) = \int p(\phi_g^j|\vartheta_j)p(\vartheta_j|\mathbb{Y})d\vartheta_j, \quad (\text{C.1})$$

where

$$p(\phi_g^j | \vartheta_j) = \sum_{c=1}^{\mathcal{K}_j} \frac{n_c^j}{\eta_j + \mathcal{N}_{\mathcal{A}}} \mathbf{N}(\phi_g^j | a_c^{j*}, h_{\phi_c}^{2j*}) + \frac{\eta_j}{\eta_j + \mathcal{N}_{\mathcal{A}}} p(\phi_g^j). \quad (\text{C.2})$$

Recall  $p(\phi_g^j)$  is the prior predictive distribution [see (B.23)]. The posterior predictive distribution (C.1) can be approximated by

$$p(\phi_g^j | \mathbb{Y}) \approx \frac{1}{R} \sum_{r=1}^R p(\phi_g^j | \vartheta_j^{(r)}). \quad (\text{C.3})$$

There is a similar expression for the posterior predictive distribution for  $\sigma_g^2$ .

### Specific fund-observation parameters

Each fund-observation, indexed by  $(i, t)$ , is associated with a regime  $s_{it} = m$ . Let  $\phi_{it}$  denote the fund-regime coefficients for fund  $i$  at time  $t$ . Referring to (B.1a), we have

$$p(\phi_{it} | Y_{iT_i}, s_{it} = m) = \mathbf{N}(\phi_{it} | \bar{\mu}_{im}, \bar{B}_{im}). \quad (\text{C.4})$$

Then,

$$p(\phi_{it} | \mathbb{Y}) = \sum_{m=1}^{M_i} p(\phi_{it} | Y_{iT_i}, s_{it} = m) p(s_{it} = m | \mathbb{Y}). \quad (\text{C.5})$$

This distribution may be approximated as follows:

$$p(\phi_{it} | \mathbb{Y}) \approx \frac{1}{R} \sum_{r=1}^R p(\phi_{it} | Y_{iT_i}, s_{it}^{(r)}) = \frac{1}{R} \sum_{r=1}^R \mathbf{N}(\phi_{it} | \bar{\mu}_{im}^{(r)}, \bar{B}_{im}^{(r)}). \quad (\text{C.6})$$

Similarly, the marginal distribution for  $\phi_{it}^j$  may be approximated by

$$p(\phi_{it}^j | \mathbb{Y}) \approx \frac{1}{R} \sum_{r=1}^R p(\phi_{it}^j | Y_{iT_i}, s_{it}^{(r)}), \quad (\text{C.7})$$

where

$$p(\phi_{it}^j | Y_{iT_i}, s_{it}) = \int p(\phi_{it} | Y_{iT_i}, s_{it}) d\phi_{it}^{-j} = \mathbf{N}(\phi_{it}^j | \bar{\mu}_{im}^j, \bar{B}_{im}^{jj}). \quad (\text{C.8})$$

For every fund-coefficient, indexed by  $(i, j)$ , we have a sequence of posterior distributions:

$$\{p(\phi_{it}^j | \mathbb{Y})\}_{t=\tau_i}^{T_i}. \quad (\text{C.9})$$

We will examine various features of these distributions. If there are no regime changes in going from  $t$  to  $t+1$ , then the two distributions will be the same. For a fund with no regime changes, the entire sequence of distributions will be the same.

## Population distribution of the regime-switching probabilities

The conditional distribution (A.34a) delivers a uniform prior predictive for the probability of a regime change (regardless of the prior for  $k_c^*$ ):

$$p(q_i) = \text{Uniform}(0, 1). \quad (\text{C.10})$$

## Posterior

We are interested in the predictive distribution of the transition probability for the unobserved  $g$ th fund

$$p(q_g|\mathbb{Y}) = \int p(q_g|\vartheta_q) p(\vartheta_q|\mathbb{Y}) d\vartheta_q, \quad (\text{C.11})$$

where

$$p(q_g|\vartheta_q) = \sum_{c=1}^{\mathcal{K}_q} \frac{n_c^q}{\eta_q + J} \text{Beta}(q_g|j_c^*, k_c^* - j_c^* + 1) + \frac{\eta_q}{\eta_q + J}. \quad (\text{C.12})$$

Recall  $p(\phi_g^j)$  is the prior predictive distribution [see (B.23)]. The posterior predictive distribution (C.1) can be approximated by

$$p(q_g|\mathbb{Y}) \approx \frac{1}{R} \sum_{r=1}^R p(q_g|\vartheta_q^{(r)}). \quad (\text{C.13})$$

## Specific regime-switching probabilities

The posterior distribution for an observed  $i$ th fund's probability of switching a regime  $q_i$ , is given by

$$p(q_i|\mathbb{Y}) = \int p(q_i|\psi_c^{q*}) p(\psi_c^{q*}|\mathbb{Y}) d\psi_c^{q*}, \quad (\text{C.14})$$

where  $c = z_i^q$  and  $i = 1, \dots, J$ . We can approximate this posterior distribution by

$$p(q_i|\mathbb{Y}) \approx \frac{1}{R} \sum_{r=1}^R \text{Beta}(q_i|j_{c^{(r)}}^{*(r)}, k_{c^{(r)}}^{*(r)} - j_{c^{(r)}}^{*(r)} + 1), \quad (\text{C.15})$$

where  $c^{(r)} = z_i^{q^{(r)}}$ .

## Forecasting

Additional predictive states can be drawn via

$$p(s_{i,T_i+1}, \dots, s_{i,T_i+H} | s_{i,T_i}, q_i) = \prod_{h=1}^H p(s_{i,T_i+h} | s_{i,T_i+h-1}, q_i). \quad (\text{C.16})$$

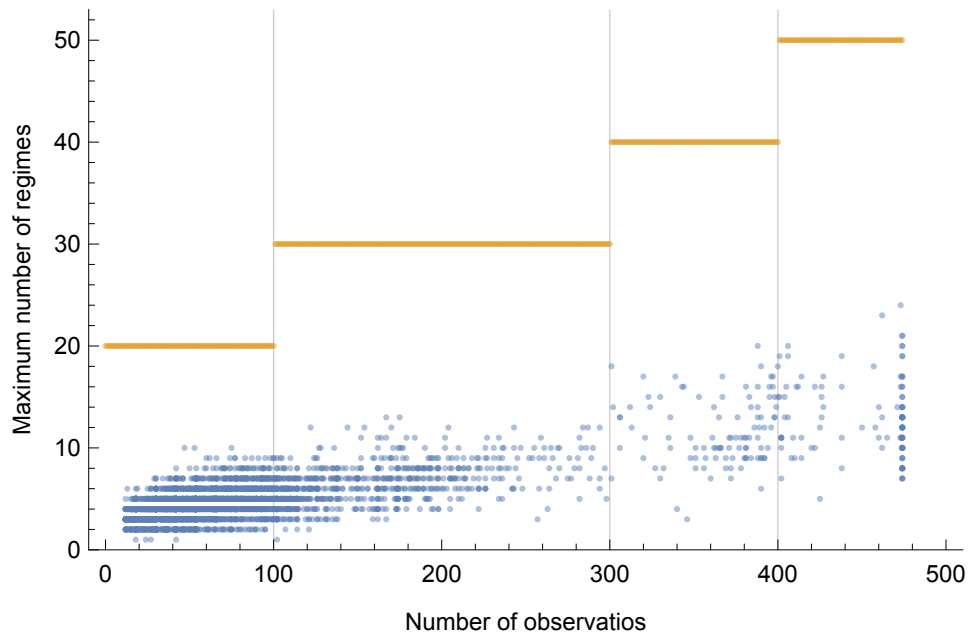


Figure 8: Maximum number of regimes versus number of observations, fund-by-fund. Maximum number allowed shown for reference.

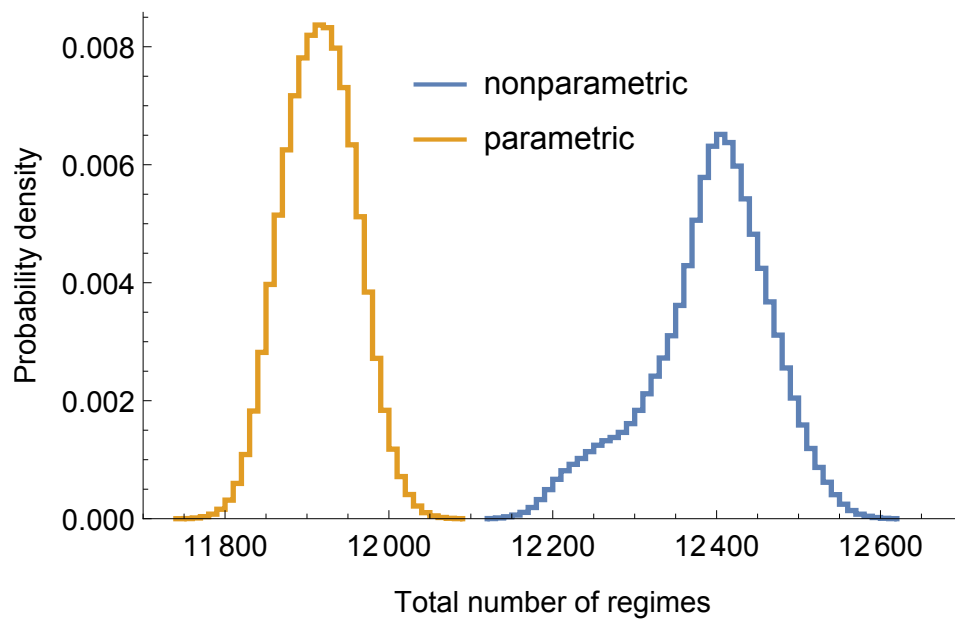


Figure 9: Smoothed histograms for total number of regimes.

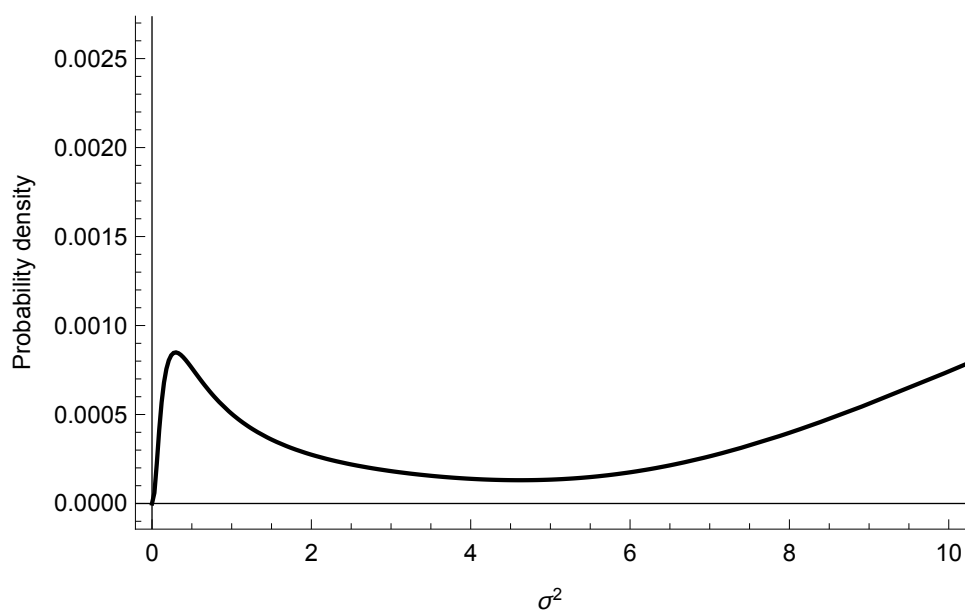


Figure 10: Detail of the posterior predictive distribution for  $\sigma^2$  showing the mode located near  $\sigma^2 = 1$ .