

ASSESSING SYSTEMATIC RISK IN THE S&P500 INDEX BETWEEN 2000 AND 2011: A BAYESIAN NONPARAMETRIC APPROACH¹

BY ABEL RODRÍGUEZ*, ZIWEI WANG[†] AND ATHANASIOS KOTTAS*

University of California, Santa Cruz and IAC Publishing Labs[†]*

We develop a Bayesian nonparametric model to assess the effect of systematic risks on multiple financial markets, and apply it to understand the behavior of the S&P500 sector indexes between January 1, 2000 and December 31, 2011. More than prediction, our main goal is to understand the evolution of systematic and idiosyncratic risks in the U.S. economy over this particular time period, leading to novel sector-specific risk indexes. To accomplish this goal, we model the appearance of extreme losses in each market using a superposition of two Poisson processes, one that corresponds to systematic risks that are shared by all sectors, and one that corresponds to the idiosyncratic risk associated with a specific sector. In order to capture changes in the risk structure over time, the intensity functions associated with each of the underlying components are modeled using a Dirichlet process mixture model. Among other interesting results, our analysis of the S&P500 index suggests that there are few idiosyncratic risks associated with the consumer staples sector, whose extreme negative log returns appear to be driven mostly by systematic risks.

1. Introduction. *Systematic risk* can be loosely defined as the vulnerability of a financial market to events that affect all (or at least most) of the agents and products in the market. This is in contrast to *idiosyncratic risks*, which are risks to which only specific agents or products are vulnerable.

The notions of systematic and idiosyncratic risks play a key role in motivating investment diversification. In this paper, we study the evolution of systematic and idiosyncratic risks in the U.S. economy by focusing on the behavior of the S&P500 index and its sector components. The Standard & Poor's 500, or S&P500 index, is a commonly watched stock market index in the U.S., constructed as a market-value weighted average of the prices of the common stock of 500 publicly traded companies. Standard & Poor's, which publishes the index, selects the companies included in the S&P500 index to be representative of the industries in the U.S. economy. These companies are commonly grouped into ten economic sectors:

Received April 2014; revised August 2016.

¹Supported in part by the NSF under awards SES 1024484 and DMS 0915272, and by funds granted by SIGFIRM at the University of California, Santa Cruz.

Key words and phrases. Dirichlet process mixture modeling, nonhomogeneous Poisson process, nonparametric Bayes, systematic risk.

consumer discretionary, consumer staples, energy, financials, health care, industrials, materials, information technology, telecommunication services and utilities. The largest sector (consumer discretionary) includes 81 companies, whereas the smallest (telecommunication services) includes only 8. In addition to the overall S&P500 index, Standard & Poor's publishes separate indexes for each of these sectors. The behavior of these sector-specific indexes is of independent interest; for example, the performance of the industrial and consumer discretionary components of the S&P500 is used by some analysts as a leading indicator of future economic growth.

The most widely used tool to characterize idiosyncratic risks in financial markets is the Capital Asset Pricing Model (CAPM) [French (2003), Treynor (1961, 1962)]. The form of the CAPM can be derived from a structural model in which agents maximize the expected utility derived from the investment, which is a function of the expected return on risky activities and the associated variance, as well as those of the market as a whole. The parameters can then be estimated using linear regression to relate the expected returns of an individual security or sector to those of the market. The estimated regression coefficient (the so-called "beta") measures the sensitivity of the expected excess asset returns to the expected excess market returns, with larger values of "beta" indicating investments with substantial idiosyncratic risks.

The original CAPM has been repeatedly criticized as being too simplistic [e.g., Fama and French (2004)], and extensions have been driven by considerations both empirical [e.g., the three factor model discussed in Fama and French (1992)] and theoretical [e.g., the behavioral model discussed in Daniel, Hirshleifer and Subrahmanyam (2001)]. Two obvious concerns with estimates of idiosyncratic risk based on the CAPM are the assumption that deviations from the model follow Gaussian distributions and their reliance on expected returns. Indeed, from a risk management perspective, it is more appealing to define these concepts on the basis of the behavior of extreme returns. As the financial crises of 2007 and 2008 demonstrated, the behavior of markets (e.g., the level of correlation among asset prices) during periods of distress can dramatically deviate from their behavior during periods of calm. Variations of CAPM that focus on extreme returns include Barnes and Hughes (2002), Allen et al. (2009) and Chang, Hung and Nieh (2011). In these papers, quantile regression instead of ordinary linear regression is used to relate the returns of individual securities to those of the market.

This paper develops a novel approach to estimate systematic and idiosyncratic market risks. Unlike the CAPM model, our focus is on tail risks, that is, risks associated with rare extreme losses (e.g., those associated with negative asset price movements of more than three standard deviations), and on reduced-form models for the relative frequency of extreme losses rather than structural models for the behavior of market actors. More specifically, we model the time of appearance of extreme losses in each market using a superposition of two Poisson processes, one

that corresponds to systematic risks that are shared by all sectors, and one that corresponds to the idiosyncratic risk associated with a specific sector. In order to capture changes in the risk structure over time, the intensity functions associated with each of the underlying components are modeled using a Dirichlet process (DP) mixture model [Antoniak (1974), Escobar and West (1995)]. Hence, our model can be conceptualized as an example of a Cox process [Cox (1955)]. In contrast to the CAPM setting, the proposed methodology does not rely (implicitly or explicitly) on the assumption that returns arise from a Gaussian distribution. Furthermore, our model is dynamic in nature, allowing for the structure of the different risks to evolve over time. As with the CAPM, the main goal of our analysis is a general description of the structure of the different risks and explanation rather than prediction. For example, the modeling approach results in idiosyncratic, sector-specific risk indexes, as well as metrics that quantify the overall relative importance of idiosyncratic and systematic risks on each sector.

There has been a growing interest in recent years on identifying *systemic risks*, that is, events that can trigger a collapse in a certain industry or economy. Some recently introduced measures of systemic risk include the CoVaR of Adrian and Brunnermeier (2010), the marginal expected shortfall (MSE) of Acharya et al. (2010) and the systemic risk measure (SRISK) of Acharya, Engle and Richardson (2012). Although some of the methods are relevant, these metrics are not tailored to systematic and idiosyncratic risks.

To motivate the class of models we develop, consider the daily returns associated with the ten sectors making up the S&P500 index (see also Section 4). Figure 1 presents the most extreme negative log returns on four of those sectors between January 1, 2000 and December 31, 2011. It is clear from the figure that all sectors present an increased frequency of extreme losses around periods of distress, such as the so-called “dot com” bubble burst in March of 2000 and the climax of the financial crises in September 2008. However, it is also clear that certain features are particular to specific sectors, such as the increased number of extreme returns in the energy sector in 2004 and 2005. Furthermore, even when the frequency of losses tends to increase significantly for all markets, it is not the case that extreme losses occur in all markets on exactly the same dates.

The rest of the paper is organized as follows. In Section 2, we develop the modeling approach, and in Section 3, we discuss posterior simulation, with technical details included in the supplementary material [Rodríguez, Wang and Kottas (2017)], as well as subjective prior elicitation. Section 4 considers the analysis of the U.S. market, using data from the S&P500 index. Finally, Section 5 provides concluding remarks.

2. Modeling approach. We focus on the negative log returns for the ten S&P500 sector indexes

$$x_{i,j} = -100 \times \log\left(\frac{S_{i,j}}{S_{i-1,j}}\right),$$

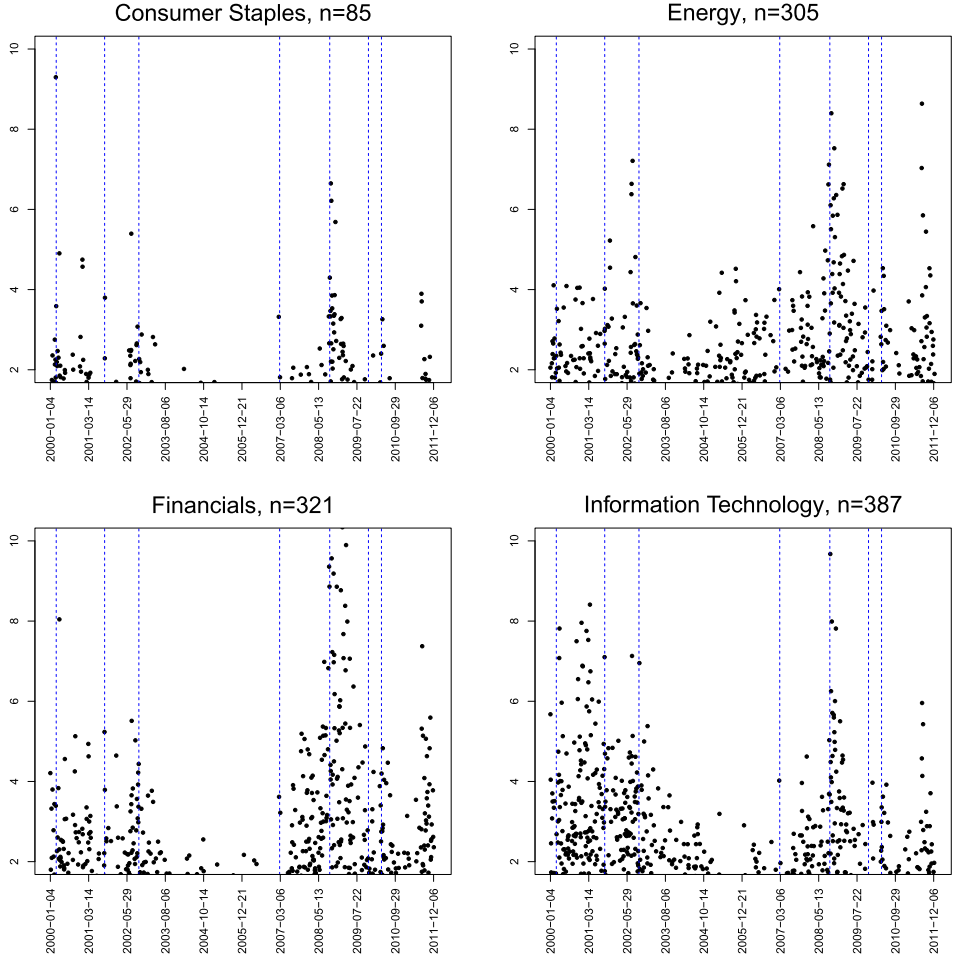


FIG. 1. Negative log returns above 2% for four sectors of the S&P 500 index (consumer staples, energy, financials and information technology). Vertical dotted lines identify seven events of significance to the markets: the bursting of the dot com bubble (03/10/2000), the 9/11 terrorist attacks (09/11/2001), the stock market downturn of 2002 (09/12/2002), the bursting of the Chinese bubble (02/27/2007), the bankruptcy of Lehman Brothers (09/16/2008), Dubai's debt standstill (11/27/2009) and the beginning of the European sovereign debt crisis (08/27/2010).

where $S_{i,j}$ is the value of the index for sector $j = 1, \dots, J = 10$ at time $i = 1, \dots, T$. Note that large positive values of $x_{i,j}$ indicate a large drop in the price index associated with sector j , and thus for risk management purposes we are interested in *large* values of $x_{i,j}$. Hence, for a given threshold u , we focus our attention on the collections of times $\{t_{j,k} : k = 1, \dots, n_j, j = 1, \dots, J\}$, where $t_{j,k}$ is the date associated with the appearance of the k th negative log return in sector j that is larger than u .

For each sector j , we regard the collection of times $\{t_{j,k} : k = 1, \dots, n_j\}$ at which exceedances occur as a realization from a point process $N_j(t)$ defined on $[0, T]$, that is, $N_j(t) = \sum_{k=1}^{n_j} \mathbb{I}_{[t_{j,k}, T]}(t)$, where $\mathbb{I}_\Omega(t)$ denotes the indicator function of set Ω . In turn, each $N_j(t)$ is constructed as the superposition of two independent, nonhomogeneous Poisson processes. The first such process accounts for systematic risk and has a cumulative intensity function Λ_0^* that is common to all sectors, while the second is associated with the idiosyncratic risk and has a cumulative intensity function Λ_j^* that is specific to each sector. Because of properties of superpositions of Poisson processes, this assumption implies that each $N_j(t)$ is also a nonhomogeneous Poisson process with cumulative intensity $\Lambda_j(t) = \Lambda_0^*(t) + \Lambda_j^*(t)$ and intensity function $\lambda_j(t) = \lambda_0^*(t) + \lambda_j^*(t)$, where λ_0^* and λ_j^* are the Poisson process intensities associated with Λ_0^* and Λ_j^* , respectively.

The modeling approach for the Λ_j^* builds from the direct connection of a nonhomogeneous Poisson process cumulative intensity/intensity function with a distribution/density function. Specifically, for $j = 0, 1, \dots, J$, we can write $\Lambda_j^*(t) = \gamma_j^* F_j^*(t)$, where $\gamma_j^* \equiv \Lambda_j^*(T) = \int_0^T \lambda_j^*(t) dt$ ($< \infty$) is the rate parameter controlling the total number of exceedances, and $F_j^*(t) = \Lambda_j^*(t)/\Lambda_j^*(T)$ is a distribution function on $[0, T]$ that controls how the exceedances are distributed over time. Hence, the sector-specific cumulative intensity function Λ_j can be written as

$$\Lambda_j(t) = \gamma_j F_j(t) = \{\gamma_0^* + \gamma_j^*\} \left\{ \frac{\gamma_0^*}{\gamma_0^* + \gamma_j^*} F_0^*(t) + \frac{\gamma_j^*}{\gamma_0^* + \gamma_j^*} F_j^*(t) \right\}.$$

This construction implies that the sector-specific exceedance rate, γ_j , is the sum of the systematic and idiosyncratic rates, while the sector-specific distribution function, F_j , can be written as a mixture of the systematic and idiosyncratic distribution functions. The corresponding weight, $\epsilon_j = \gamma_0^*/(\gamma_0^* + \gamma_j^*)$, represents the proportion of exceedances in sector j that are associated with the systematic component. In addition, note that values of ϵ_j close to 1 (which are associated with $\gamma_0^* \gg \gamma_j^*$) imply a stronger association in the pattern of extreme losses.

Because each $N_j(t)$ follows a Poisson process, the probability that at most r exceedances will be observed in sector j during time period $[t_0, t_0 + \Delta]$ is

$$\sum_{i=0}^r \frac{\{\Upsilon_j(t_0, \Delta)\}^i \exp\{-\Upsilon_j(t_0, \Delta)\}}{i!},$$

where $\Upsilon_j(t_0, \Delta) = \Lambda_j(t_0 + \Delta) - \Lambda_j(t_0)$. These exceedance probabilities are easier to interpret than the intensity functions through which the model is defined. For example, the probability that no exceedances are observed in sector j between time points t_0 and $t_0 + \Delta$ is given by

$$\begin{aligned} \exp\{-[\Lambda_j(t_0 + \Delta) - \Lambda_j(t_0)]\} &= \exp\{-[\Lambda_0^*(t_0 + \Delta) - \Lambda_0^*(t_0)]\} \\ &\quad \times \exp\{-[\Lambda_j^*(t_0 + \Delta) - \Lambda_j^*(t_0)]\}, \end{aligned}$$

where the first term in the right-hand side expression corresponds to the probability of no exceedance due to the systematic component and the second term corresponds to the probability of no exceedance due to the idiosyncratic component. In this sense, our model implies a multiplicative risk structure.

2.1. Modeling the intensity functions. To generate a flexible model that can capture changes in the pattern of extreme returns over time, we model the densities f_0^* and $f_1^*, \dots, f_J^*$ associated with the systematic and idiosyncratic distribution functions F_0^* and $F_1^*, \dots, F_J^*$ using DP mixtures. In particular,

$$f_j^*(t) \equiv f^*(t|G_j^*, \tau) = \int \psi(t|\mu, \tau) dG_j^*(\mu), \quad j = 0, 1, \dots, J,$$

where $\psi(t|\mu, \tau)$ is a kernel density on $[0, T]$ indexed by parameters μ and τ , and G_j^* is a discrete mixing distribution.

By carefully choosing the kernel ψ and the prior on the mixing distribution G_j^* , we can obtain a flexible model for the intensity functions. Hence, we opt for a prior on G_j^* that has full support on the space of discrete distributions. In particular, G_j^* is assigned a DP prior [Ferguson (1973)] with precision parameter α_j and baseline (centering) distribution H , which is common to all G_j^* . Thus, using the stick-breaking DP constructive definition [Sethuraman (1994)],

$$G_j^*(\cdot) = \sum_{l=1}^{\infty} \left\{ v_{j,l} \prod_{s<l} (1 - v_{j,s}) \right\} \delta_{\tilde{\mu}_{j,l}}(\cdot), \quad j = 0, 1, \dots, J,$$

where $\delta_a(\cdot)$ denotes the degenerate measure at a , the atoms $\{\tilde{\mu}_{j,1}, \tilde{\mu}_{j,2}, \dots\}$ form a sequence of random variables independent and identically distributed according to H , and $\{v_{j,1}, v_{j,2}, \dots\}$ is another sequence (independent of $\{\tilde{\mu}_{j,1}, \tilde{\mu}_{j,2}, \dots\}$) of independent and identically distributed random variables according to a $\text{Beta}(1, \alpha_j)$ distribution.

Because in our application the support for the point process is a compact set, a natural choice for the kernel $\psi(t|\mu, \tau)$ is the rescaled beta density,

$$(2.1) \quad \frac{1}{T} \frac{\Gamma(\tau)}{\Gamma(\mu\tau/T)\Gamma(\{1 - \mu/T\}\tau)} \left(\frac{t}{T}\right)^{\mu\tau/T-1} \left(1 - \frac{t}{T}\right)^{\{1 - \mu/T\}\tau-1} \mathbb{I}_{[0,T]}(t),$$

where $\mu \in [0, T]$ is a location parameter (the mean of the kernel distribution), and $\tau > 0$ can be interpreted as a scale parameter.

Note that if $\mu\tau/T > 1$ and $\{1 - \mu/T\}\tau > 1$, then the rescaled beta density is unimodal, with the mode located approximately at μ , for large τ . Hence, the location mixture of beta kernels can accommodate the type of risk clustering we observe in Figure 1. Furthermore, because DP mixtures allow for a countable number of mixture components, the model is dense on the space of absolutely continuous distributions on $[0, T]$ as long as the baseline distribution H and the prior on the scale parameter τ are selected to provide full support on the domain of the parameters of the rescaled beta kernel [see, e.g., Diaconis and Ylvisaker (1985)]. Note

that the precision parameter α_j controls the relative weight of the mixture components, with smaller values of α_j favoring mixtures with a small number of effective components. On the other hand, the baseline distribution H controls the location of the mixture components.

Besides a prior on the densities $f_0^*, f_1^*, \dots, f_J^*$, full prior specification for the intensity functions $\lambda_1(t), \dots, \lambda_J(t)$ requires priors for the rate parameters $\gamma_0^*, \gamma_1^*, \dots, \gamma_J^*$. In the case of the rate associated with the common component, γ_0^* , a natural choice is a gamma distribution with shape parameter $a_{\gamma_0^*}$ and rate parameter $b_{\gamma_0^*}$. For the idiosyncratic component, we use a more general, zero-inflated gamma prior with density,

$$p(\gamma_j^*|\pi) = (1 - \pi)\delta_0(\gamma_j^*) + \pi \text{Gam}(\gamma_j^*|a_{\gamma_j^*}, b_{\gamma_j^*}), \quad j = 1, \dots, J,$$

where $\text{Gam}(a, b)$ denotes a gamma distribution with mean a/b . Note that the case $\gamma_j^* = 0$ corresponds to $\epsilon_j = 1$, that is, all exceedances in sector j are driven by systematic risks. Hence, this zero-inflated prior allows us to formally test for the presence of idiosyncratic risks. In the sequel we refer to this test as the *idiosyncrasy test*.

2.2. Hierarchical priors. The model is completed by adding hyperpriors to model parameters. For the baseline distribution H , we note that the parameter μ in equation (2.1) is constrained to the $[0, T]$ interval. Hence, a natural choice for H is another rescaled beta distribution with density $h(\mu) \propto (\mu/T)^{a_\mu-1} \{1 - (\mu/T)\}^{b_\mu-1} \mathbb{I}_{[0, T]}(\mu)$. Although the full Bayesian model can be extended to incorporate a hyperprior for one or both of the parameters of H , this was not necessary for the application to the S&P500 data. In fact, we applied the model using the default choice of a uniform distribution for H ; sensitivity of posterior inference results to this choice, as well as to the other hyperprior choices, is discussed in Section 4.2.

The remaining priors are selected for computational convenience. In particular, a conditionally conjugate $\text{Gam}(a_\alpha, b_\alpha)$ prior is assigned to the DP precision parameters α_j , $j = 0, 1, \dots, J$. Similarly, the conditionally conjugate prior for the mixing probability π is given by a $\text{Beta}(a_\pi, b_\pi)$ distribution. Finally, the scale parameter τ is assigned an inverse gamma distribution with shape a_τ and scale b_τ [denoted by $\text{IGam}(a_\tau, b_\tau)$]. Details on informative hyperparameter elicitation are discussed in Section 3.2

2.3. Related literature. The DP mixture modeling approach for the systematic and idiosyncratic intensities builds on earlier work for temporal or spatial Poisson processes [Kottas and Sansó (2007), Taddy and Kottas (2012)]. This methodology has been applied to inference for neuronal firing intensities [Kottas et al. (2012)], tracking the intensity of violent crime [Taddy (2010)], risk assessment for extremes of environmental time series [Kottas, Wang and Rodríguez (2012)] and analysis of

hurricane landfall occurrences [Xiao, Kottas and Sansó (2015)]. These applications involve modeling of dependent Poisson process intensities, with the index of dependence ranging from a finite set indicating different experimental conditions [Kottas et al. (2012)], to discrete time [Taddy (2010), Xiao, Kottas and Sansó (2015)], to space [Kottas, Wang and Rodríguez (2012)]. Dependence in the prior model for the intensities is built through dependent DP [MacEachern (2000)] or spatial DP [Gelfand, Kottas and MacEachern (2005)] priors for the corresponding mixing distributions that define the mixture model for the Poisson process densities. Bassetti, Casarin and Leisen (2014) discuss other prior models for finite collections of distributions, including dependent Pitman–Yor process priors which extend dependent DP priors.

This paper provides an example of a more structured dependent prior model for a finite collection of Poisson process intensities. The representation of the sector-specific distribution function F_j as a mixture of a systematic and an idiosyncratic component is reminiscent of the models for dependent random distributions discussed in Müller, Quintana and Rosner (2004), Hatjispyros, Nicolieris and Walker (2011), and Kolossiatis, Griffin and Steel (2013). Müller, Quintana and Rosner (2004) focused on the case $F_j = \varepsilon H_0 + (1 - \varepsilon)H_j$, for $j = 1, \dots, J$, where the H_j , $j = 0, 1, \dots, J$, are assigned independent DP priors with precision parameters α_j and common centering distribution. Kolossiatis, Griffin and Steel (2013) studied the more general version, $F_j = \varepsilon_j H_0 + (1 - \varepsilon_j)H_j$, albeit for two distributions, arguing for a dependent prior on $(\varepsilon_1, \varepsilon_2)$ built from $\varepsilon_j = \gamma_0/(\gamma_0 + \gamma_j)$, with independent gamma priors for γ_0 , γ_1 and γ_2 . Hatjispyros, Nicolieris and Walker (2011) used a similar modeling approach, but with a priori independent ε_1 and ε_2 . Although there is an apparent connection in model structure with these earlier methods, the motivation for our modeling approach is different. Indeed, the objective of the earlier work is to construct dependent random densities through mixture modeling, where the prior model for the F_j is used for the dependent mixing distributions. However, our model structure for the sector-specific Poisson process densities, $f_j = \epsilon_j f_0^* + (1 - \epsilon_j)f_j^*$, for $j = 1, \dots, J$, follows from the assumption that the point process associated with the presence of extreme values can be constructed as a superposition of Poisson processes. In particular, the structure with sector-specific weights of the form $\epsilon_j = \gamma_0^*/(\gamma_0^* + \gamma_j^*)$ is a direct consequence from the superposition of the Poisson processes, rather than a particular model choice as in Kolossiatis, Griffin and Steel (2013).

Regarding the point process literature, a relevant reference is Griffiths and Milne (1978), where a bivariate point process was studied such that the marginals, $N_1(t)$ and $N_2(t)$, follow the same Poisson process with cumulative intensity $\Lambda(\cdot)$. The key result provides a specific form for the Laplace transform of the bivariate point process which is achieved if and only if the marginals admit the representation, $N_1(t) = M_1(t) + M_2(t)$ and $N_2(t) = M_2(t) + M_3(t)$, where $M_1(t)$, $M_2(t)$ and $M_3(t)$ are independent Cox processes with respective cumulative intensities $\nu(\cdot)$, $\mu(\cdot)$ and $\nu(\cdot)$, with $\mu(\cdot) \leq \Lambda(\cdot)$ almost surely, and $\nu(\cdot) = \Lambda(\cdot) - \mu(\cdot)$. The Laplace

transform is the main functional used in the study of completely random measures (CRMs), which are defined in terms of underlying Poisson processes and have been used as general nonparametric prior models for distributions; see, for example, Lijoi and Prünster (2010) for a review. Hence, this result has inspired work on dependent CRMs, built from dependent Poisson processes that admit the representation above. We refer to Lijoi, Nipoti and Prünster (2014a, 2014b) for related theory and inference methods in the context of mixture models, as well as to Lijoi and Nipoti (2014) for modeling dependent hazard rate functions in survival analysis. [Alternative methods for construction of dependent CRMs include Leisen and Lijoi (2011) and Leisen, Lijoi and Spanó (2013).] Although this work is based on the same representation for Poisson process intensities that we use in our approach, the objective is to build nonparametric priors for dependent distributions. In particular, the representation in terms of a common term and an idiosyncratic term is utilized mainly as a mechanism to build dependence in the prior model for two correlated distributions. In our approach, the particular representation is the key modeling assumption for the occurrence of extreme values, and inference for the systematic and idiosyncratic components is of direct interest.

3. Posterior simulation and prior elicitation.

3.1. Markov chain Monte Carlo posterior inference. Our modeling approach is based on the decomposition of the Poisson process intensity into the total intensity, $\gamma_j = \gamma_0^* + \gamma_j^*$, and the Poisson process density, $f_j(t) = \epsilon_j f_0^*(t) + (1 - \epsilon_j) f_j^*(t)$, for $t \in [0, T]$. Under this decomposition, the likelihood function associated with the nonhomogeneous Poisson process giving rise to the exceedances in sector j is $\gamma_j^{n_j} \exp\{-\gamma_j\} \prod_{k=1}^{n_j} f_j(t_{j,k})$. Hence, the joint posterior distribution for our model can be written as

$$\begin{aligned}
 & p(\{\gamma_j^*\}, \{v_{j,l}\}, \{\tilde{\mu}_{j,l}\}, \tau, \{\alpha_j\}, \pi | \text{data}) \\
 & \propto \prod_{j=1}^J (\gamma_0^* + \gamma_j^*)^{n_j} \exp\{-(\gamma_0^* + \gamma_j^*)\} \\
 (3.1) \quad & \times \prod_{j=1}^J \prod_{k=1}^{n_j} \left(\frac{\gamma_0^*}{\gamma_0^* + \gamma_j^*} \sum_{l=1}^{\infty} \left\{ v_{0,l} \prod_{s < l} (1 - v_{0,s}) \right\} \psi(t_{j,k} | \tilde{\mu}_{0,l}, \tau) \right. \\
 & \left. + \frac{\gamma_j^*}{\gamma_0^* + \gamma_j^*} \sum_{l=1}^{\infty} \left\{ v_{j,l} \prod_{s < l} (1 - v_{j,s}) \right\} \psi(t_{j,k} | \tilde{\mu}_{j,l}, \tau) \right) \\
 & \times p(\pi) p(\tau) p(\gamma_0^*) \prod_{j=1}^J p(\gamma_j^* | \pi) \prod_{j=0}^J \prod_{l=1}^{\infty} h(\tilde{\mu}_{j,l}) \text{Beta}(v_{j,l} | 1, \alpha_j) \prod_{j=0}^J p(\alpha_j),
 \end{aligned}$$

where the priors $p(\pi)$, $p(\tau)$, $p(\gamma_0^*)$, $p(\gamma_j^*|\pi)$, for $j = 1, \dots, J$, and $p(\alpha_j)$, for $j = 0, 1, \dots, J$, are given in Section 2.

Since the posterior distribution is computationally intractable, we resort to a Markov chain Monte Carlo (MCMC) algorithm [Robert and Casella (2004)] for simulation-based inference. To sample from the posterior distribution associated with the nonparametric component of the model, we use blocked Gibbs sampling [Ishwaran and James (2001), Ishwaran and Zarepour (2000)]. Hence, for computational purposes, we replace the mixing distributions G_0^* , G_1^* , $\dots$, G_J^* with finite-dimensional approximations:

$$G_j^N(\cdot) = \sum_{l=1}^N \left\{ u_{j,l} \prod_{s < l} (1 - u_{j,s}) \right\} \delta_{\mu_{j,l}^*}(\cdot), \quad j = 0, 1, \dots, J,$$

where, as before, the $\mu_{j,l}^*$ are drawn independently from H and the $u_{j,l}$ are independent $\text{Beta}(1, \alpha_j)$, for $l = 1, \dots, N-1$, but $u_{j,N} = 1$ to ensure that the weights sum to 1. The truncation level N can be chosen to any desired level of accuracy using distributional results for the DP stick-breaking weights. In particular, $E(\sum_{l=1}^N u_{j,l} \prod_{s < l} (1 - u_{j,s}) | \alpha_j) = 1 - \{\alpha_j / (\alpha_j + 1)\}^N$, which can be averaged over the prior for α_j to estimate the prior expectation for the partial sum of the first N weights. In practice, in addition to using the above result to guide the selection of N , we perform a sensitivity analysis to investigate the impact of our choice (see Section 4).

Furthermore, we expand the model by introducing for each observation $t_{j,k}$ a pair of latent configuration variables, $(r_{j,k}, L_{j,k})$, where $r_{j,k} \in \{0, j\}$ is an indicator for the systematic or idiosyncratic component, whereas $L_{j,k} \in \{1, \dots, N\}$ identifies the respective mixture component under the truncation approximation for G_0^N or G_j^N . More specifically, independently for $j = 1, \dots, J$ and $k = 1, \dots, n_j$, $\Pr(r_{j,k} = 0 | \gamma_0^*, \gamma_j^*) = 1 - \Pr(r_{j,k} = j | \gamma_0^*, \gamma_j^*) = \gamma_0^* / (\gamma_0^* + \gamma_j^*)$, and, for $l = 1, \dots, N$, $\Pr(L_{j,k} = l | r_{j,k} = 0, G_0^N) = u_{0,l} \prod_{s < l} (1 - u_{0,s})$ and $\Pr(L_{j,k} = l | r_{j,k} = j, G_j^N) = u_{j,l} \prod_{s < l} (1 - u_{j,s})$. In addition to the aforementioned indicator variables associated with the mixture representation of the intensity functions, we introduce a set of binary indicators, $\xi_1, \dots, \xi_J$, such that $\Pr(\xi_j = 0 | \pi) = 1 - \pi$. Inferences on $\Pr(\xi_j = 0 | \text{data})$ provide an operational mechanism to implement the idiosyncrasy test discussed at the end of Section 2.1. Details of the posterior simulation algorithm, and on point and interval estimation for the intensity functions $\lambda_0^*, \dots, \lambda_J^*$, are provided in the supplementary material [Rodríguez, Wang and Kottas (2017)].

As pointed out by one of the referees, a possible alternative to the blocked Gibbs sampler is a slice sampler [e.g., Griffin and Walker (2011), Kalli, Griffin and Walker (2011)]. Such an approach allows posterior simulation from most model parameters without explicitly truncating the countable representation for the G_j^* , $j = 0, 1, \dots, J$, and is therefore attractive if the inference objectives are limited to posterior predictive estimation which does not require posterior samples for the

mixing distributions. However, key for our analysis is inference for the sector-specific intensities, $\lambda_j(t)$, $j = 1, \dots, J$, defined through mixture densities $f_j^*(t)$, $j = 0, 1, \dots, J$, and for the probability of r exceedances over a given period. And, we are interested not only in point estimates for these functionals, but also in estimates of the uncertainty. Obtaining those estimates requires posterior samples for the mixing distributions, for which truncation is necessary also under a slice sampler. Hence, the advantages of slice sampling in this context are limited.

3.2. Hyperparameter elicitation. To elicit the model hyperparameters, historical and/or expert information can be used. Such information is typically available for most liquid financial markets. We recommend that this elicitation process be complemented with a careful sensitivity analysis over a reasonable range of prior beliefs.

Consider first the parameters $\gamma_0^*, \gamma_1^*, \dots, \gamma_{10}^*$, which control the total number of exceedances observed in each sector and the relative distribution of these exceedances between the systematic and idiosyncratic component of the model. Because of their role in the model, we can elicit a value for the expected number of extreme returns in a given sector j [which corresponds to $E(\gamma_0^* + \gamma_j^*)$] by assuming that returns in the sector are normally distributed so that $E(\gamma_0^* + \gamma_j^*) \approx T \times \Phi(\{-u - \zeta_j\}/\kappa_j)$, where Φ denotes the standard normal distribution function, and ζ_j and κ_j are rough estimates of the mean and standard deviation of returns for sector j . The values of ζ_j and κ_j can be obtained from historical data or expert knowledge. For simplicity, it can be assumed that ζ_j and κ_j are the same for every sector, leading to a model where sectors are exchangeable, but this is not required. Similarly, we can exploit the interpretation of $\gamma_0^*/(\gamma_0^* + \gamma_j^*)$ as the proportion of exceedances arising from the systematic component to elicit expert information about the most likely value of such rate, as well as a high probability range for its value. This same information can be used to provide informative priors for $1 - \pi$, the prior probability that the risk in a given sector is entirely driven by the systematic component.

Next, consider eliciting the hyperparameters associated with the densities f_0^* and $f_1^*, \dots, f_J^*$. A common feature of extreme returns in financial time series is that they tend to cluster over time [e.g., Mandelbrot (1963)]. Hence, the priors for the precision parameters $\alpha_0, \alpha_1, \dots, \alpha_J$ (which, as mentioned in Section 2, control the number of mixture components) should support relatively large values. A rough value for the number of components, which can be used to select the prior mean of α_j [see, e.g., Escobar and West (1995)], can be elicited from a rough estimate of the frequency at which distress periods arise in sector j . Similarly, the value for the scale parameter τ can be elicited from prior information about the length of distress periods. Finally, in the absence of prior information about the time at which distress periods occur, we recommend that H be selected so that the prior mean for f_0^* and $f_1^*, \dots, f_J^*$ is close to uniform.

TABLE 1
S&P500 sector indexes and their associated tickers

Sector number	Ticker	Sector description
1	S5COND	Consumer discretionary
2	S5CONS	Consumer staples
3	S5ENRS	Energy
4	S5FINL	Financials
5	S5HLTH	Health care
6	S5INDU	Industrials
7	S5MATR	Materials
8	S5INFT	Information technology
9	S5TELS	Telecommunications
10	S5UTIL	Utilities

4. Systematic and idiosyncratic risks in the U.S. economy: An analysis of the S&P500 index and its components. The data analyzed in this section corresponds to the negative daily log returns above $u = 2\%$ on each of the ten sectors that make up the S&P500 index from the time period between January 1, 2000 and December 31, 2011. This results in sample sizes for extreme returns that range from 85 (for consumer staples) to 387 (for information technology). Prices for the individual indexes were obtained from Bloomberg financial services; see Table 1 for the corresponding tickers. A 2% drop in the market has been historically used as a threshold for trading curbs on program trades (which involve a basket of stocks from the S&P500 index where there are at least 15 stocks or where the value of the basket is at least \$1 million). Although these particular trading curbs were abandoned in late 2007 because of their inefficacy in reducing market volatility, we believe that the 2% threshold is still a useful guideline to identify periods of distress in the market without excessively thinning the sample. We repeated our analysis for other values of the threshold u without significant qualitative changes in the conclusions; refer to the supplementary material [Rodríguez, Wang and Kottas (2017)].

All inferences reported here are based on 3000 posterior samples obtained after a burn-in period of 20,000 iterations and thinning of the original chain every 50 iterations. Convergence of the MCMC algorithm was monitored using trace plots as well as the R statistic discussed in Gelman and Rubin (1992). In particular, we ran four independent chains started from overdispersed initial values and compared between and within chain variability in the value of the log-likelihood function and in some of the hyperparameters in the model. No lack of convergence was evident from these diagnostics. The algorithm was implemented in C/C++, and total execution time was approximately 16 hours on a MacBook laptop with a 2 GHz Intel Core 2 Duo processor and 2GB of memory.

Using the guidelines discussed in Section 3.1, the truncation level for the blocked Gibbs sampler was set to $N = 60$. A sensitivity analysis (using both

TABLE 2
*Parameters of the prior distributions used in the analysis of the
 S&P500 sectors*

Parameter	Prior parameters
π	$a_\pi = 0.5, b_\pi = 2$
γ_0^*	$a_{\gamma_0^*} = 7.32, b_{\gamma_0^*} = 0.06$
$\gamma_j^*, j = 1, \dots, 10$	$a_{\gamma_j^*} = 1.32, b_{\gamma_j^*} = 0.06$
τ	$a_\tau = 5, b_\tau = 2400$
$\alpha_j, j = 0, 1, \dots, 10$	$a_\alpha = 4, b_\alpha = 1/3$
H	$a_\mu = 1, b_\mu = 1$

$N = 50$ and $N = 70$) suggests that values larger than 60 lead to essentially identical results.

Table 2 summarizes our choice of prior hyperparameters. For $\gamma_0^*, \dots, \gamma_{10}^*$ and π , the choice reflects the prior belief that the systematic component of the risk explains the majority of exceedances observed in the data. Indeed, we construct our prior for π by assuming that $E(\pi) = 0.2$ and placing high probability on values of π close to zero. Moreover, we specify the priors for $\gamma_0^*, \dots, \gamma_{10}^*$ on the basis of a 0 mean return with 18% annualized volatility for the S&P500, along with the prior beliefs that on average 85% of the observed exceedances, and with 0.99 probability at least 50% of them, arise from the systematic component of the model (recall the discussion in Section 3.2). On the other hand, our choice of priors for H , for τ and for the $\alpha_j, j = 0, 1, \dots, J$, is guided by the fact that we expect highly multimodal intensity functions, but have little information about where the modes lie. Hence, our prior supports relatively large values of α_j [$E(\alpha_j) = 12$ a priori] so that individual realizations from the prior are multimodal, but we set H to the uniform distribution and favor relatively large values of τ so that the prior mean shape of the intensity function is relatively flat (see Figure 3). In all cases, the posterior distributions for all hyperparameters appear to be concentrated relative to the corresponding prior distributions. Results from prior sensitivity analysis are discussed in Section 4.2.

Estimates of the overall intensities $\lambda_1(t), \dots, \lambda_{10}(t)$ associated with each of the ten components of the S&P500 index can be seen in Figure 2. The last two panels also provide summaries of the prior distribution over intensities induced by the prior choices discussed above. By comparing some of those estimates to the raw data presented in Figure 1, it becomes clear that the model faithfully reproduces the main features of the data. Furthermore, the uncertainty associated with these estimates is relatively low. The ratio of these intensities with respect to the intensity associated with the systematic component, $\lambda_1(t)/\lambda_0^*(t), \dots, \lambda_{10}(t)/\lambda_0^*(t)$, can be used as a time-dependent index of idiosyncratic risk on each of the sectors of the U.S. economy tracked by the S&P500 index. Indeed, note that $\lambda_j(t)/\lambda_0^*(t) =$

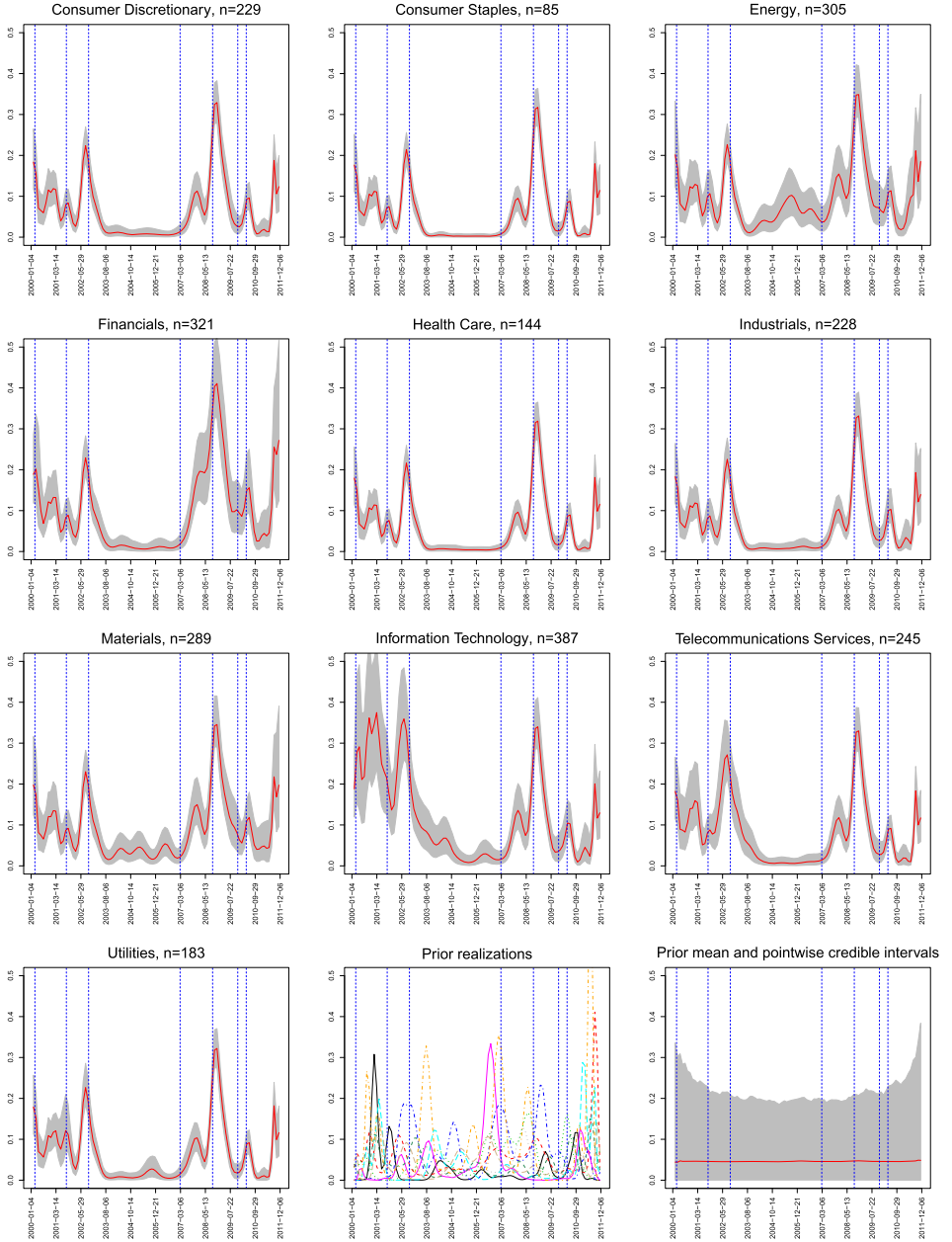


FIG. 2. The top three rows and the left panel on the bottom row show the posterior mean of the overall intensity, λ_j , $j = 1, \dots, 10$, associated with the different components of the S&P500 index, along with posterior 95% pointwise credible intervals. The headers on each panel include the number of exceedances observed in each sector over the 12-year period under study. The last two plots in the bottom row present prior realizations for the intensity function (middle panel) and the mean prior intensity function along with prior 95% pointwise credible intervals (right panel).

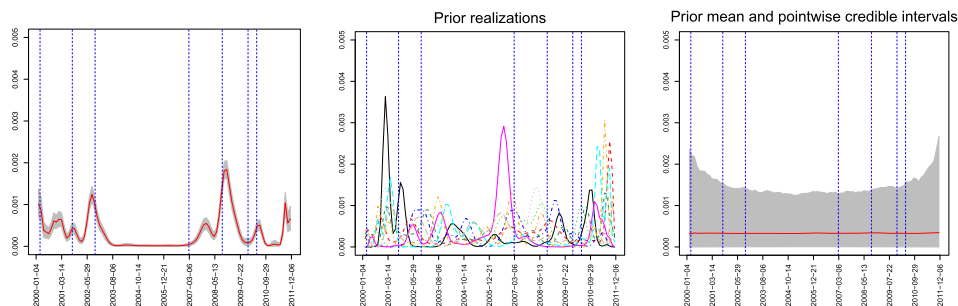


FIG. 3. The left panel shows the posterior mean of f_0^* , the density associated with the systematic risk component of the S&P500 index, including posterior 95% pointwise credible intervals. The middle panel shows prior realizations for f_0^* , while the right panel plots the prior mean and prior 95% pointwise credible intervals for f_0^* .

$1 + \{\gamma_j^* f_j^*(t) / (\gamma_0^* f_0^*(t))\}$ indicates the excess risk of a given sector relative to the baseline at time t , with sectors that exhibit little idiosyncratic risk across time having $\lambda_j(t) / \lambda_0^*(t) \approx 1$ for all t .

Next, Figures 3 and 4 show estimates of the densities f_0^* and $f_1^*, \dots, f_{10}^*$ associated with the systematic and idiosyncratic risk intensities. Figure 5 presents the posterior densities for $\epsilon_1, \dots, \epsilon_{10}$, the (average) proportion of the risk attributable to the systematic component in each of the ten sectors. Note that in half the sectors (consumer discretionary, consumer staples, health care, industrials and utilities) the proportion of extremes associated with the systematic component is at least 80%, while for the rest (energy, financials, information technology, telecommunications and materials) the proportion is between 40% and 80%. In addition, note that the density for the systematic risk shows peaks that coincide, or shortly follow, important stock market events. On the other hand, the behavior of the idiosyncratic risk varies drastically with the economic sector and, in most cases, can be explained by factors that are sector specific. For example, the energy and utilities sectors present increases in idiosyncratic risk during 2005, a period that corresponded to sharp increases in oil prices but that was otherwise relatively calm. On the other hand, the idiosyncratic risk associated with the financial sector increases dramatically after the summer of 2007. An oversized idiosyncratic risk for this sector after 2007 is clearly reasonable as financials were the main driver of the recent crisis. Similarly, the idiosyncratic risks associated with the information technology and telecommunication services sectors are particularly elevated between 2000 and 2002, a period that included the bursting of the so-called dot-com bubble. Finally, note that the idiosyncratic risk associated with consumer staples is almost negligible over the whole period under study, with our idiosyncrasy test suggesting that there is moderate evidence for the absence of idiosyncratic risk in this sector of the S&P500 index. This is reasonable, as the consumer staples sector includes companies that produce and trade basic necessities whose consumption might be affected by general economic conditions but is otherwise relatively stable.

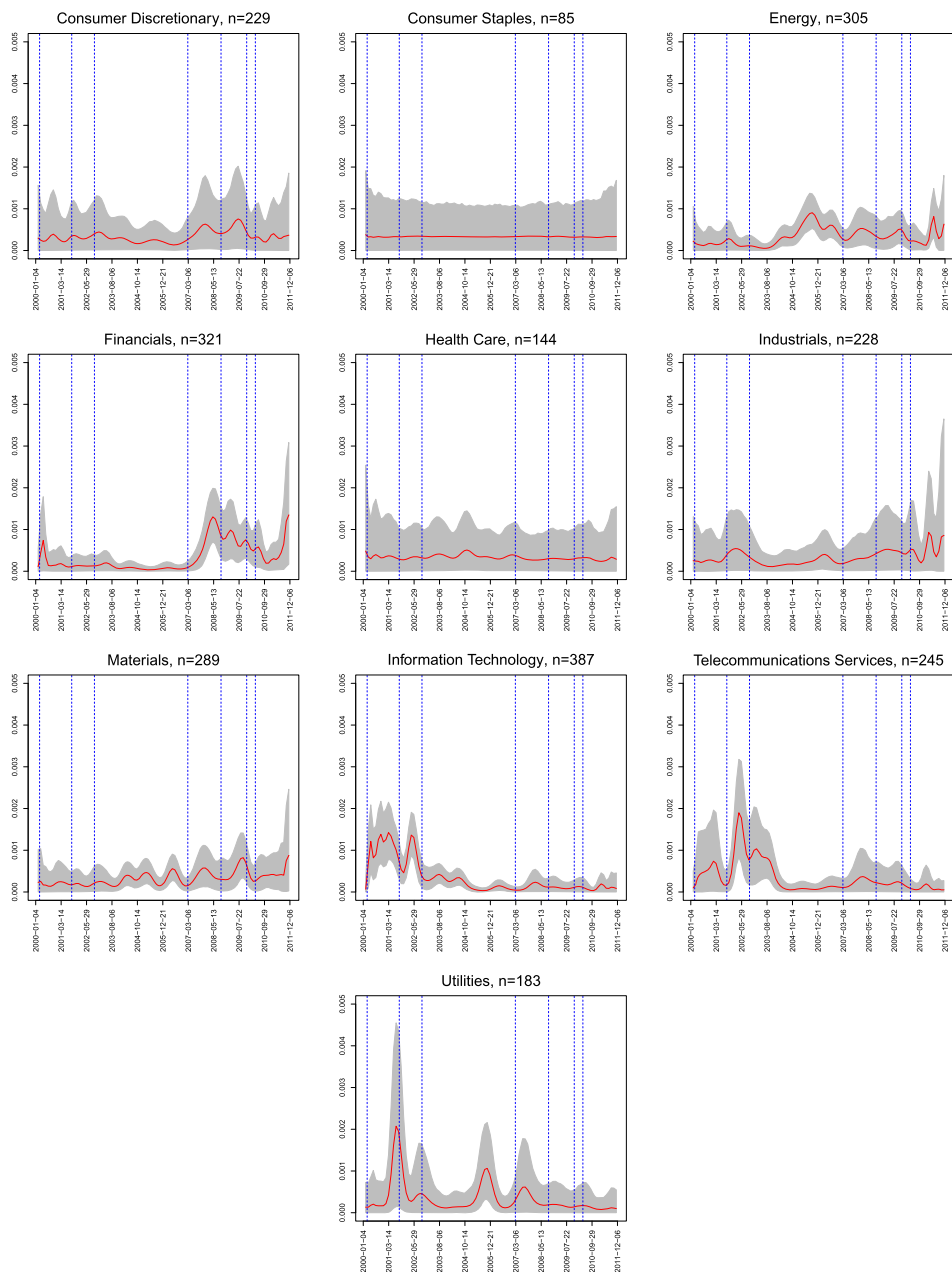


FIG. 4. Posterior mean of the idiosyncratic densities, $f_1^*, \dots, f_{10}^*$, associated with the different components of the S&P500 index, along with posterior 95% pointwise credible intervals.

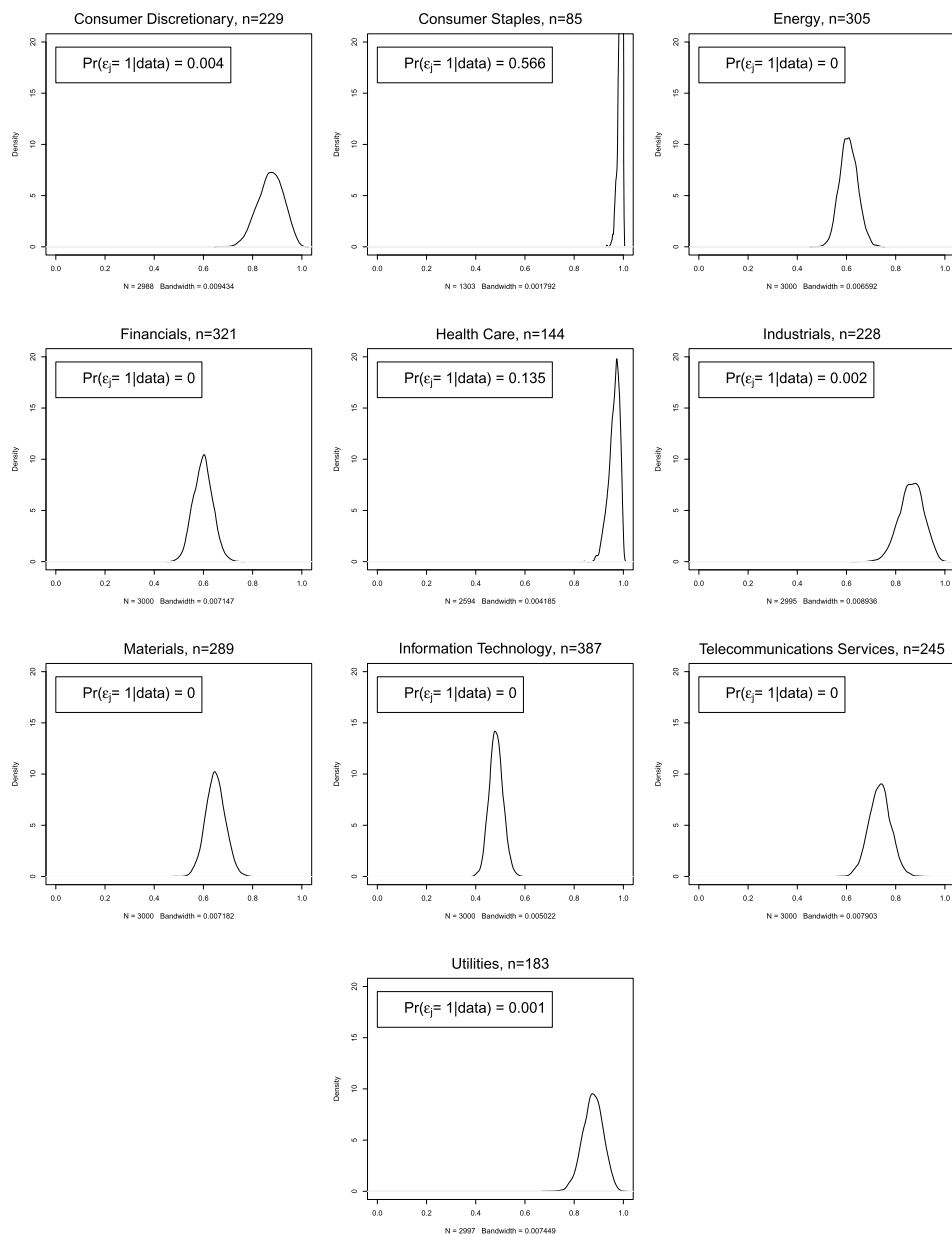


FIG. 5. Posterior density for ϵ_j , $j = 1, \dots, 10$, the overall proportion of risk attributable to the systematic component on each of the ten components of the S&P500 index.

As we discussed in Section 2, we can alternatively quantify the level of risk through the probability of observing at least one exceedance during a given period of time. Figure 6 shows for four different sectors the posterior distributions for

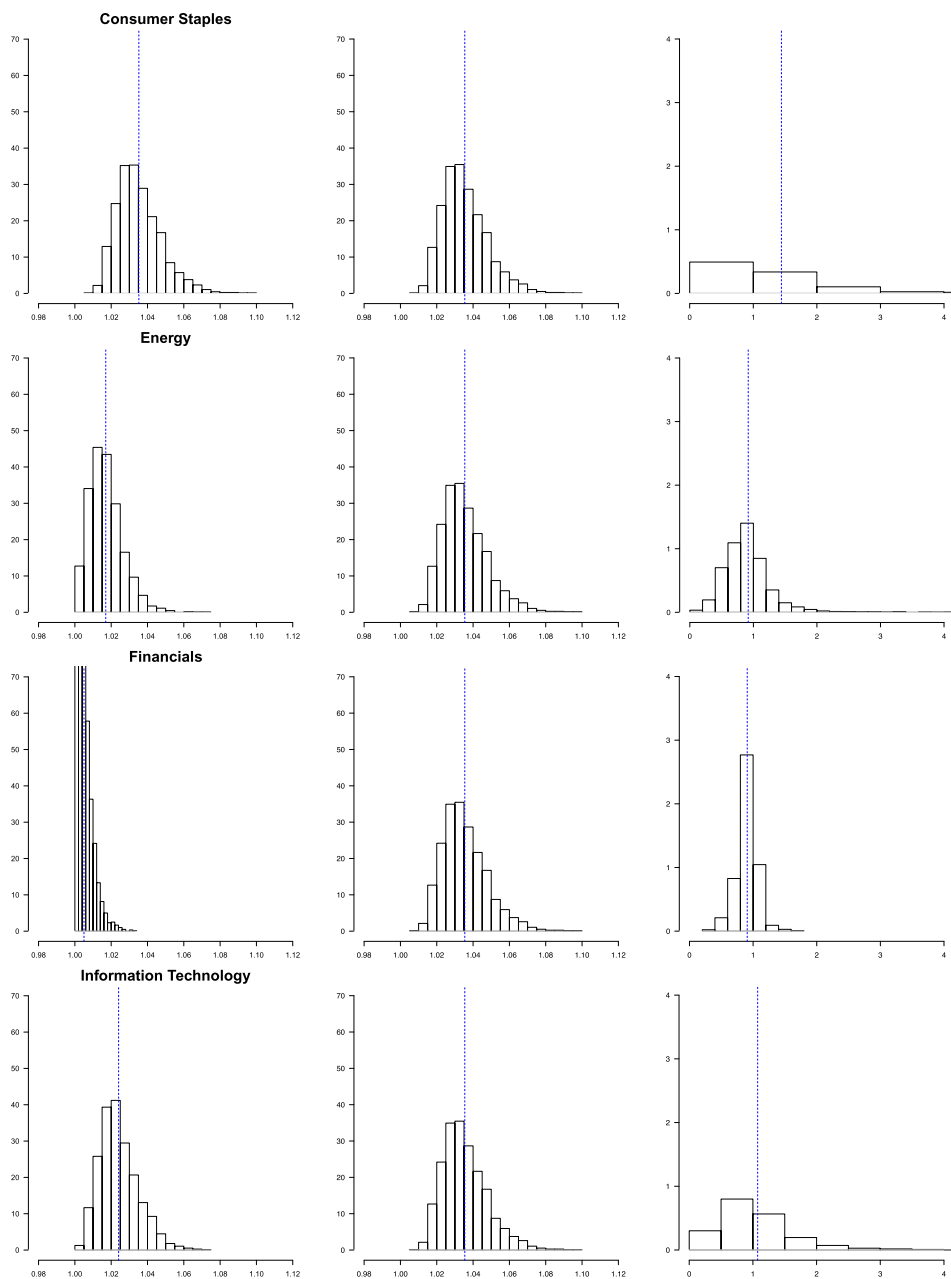


FIG. 6. The left column plots show for four different sectors the posterior densities for the odds ratio of the probability of at least one exceedance in the month starting two weeks after the bankruptcy of Lehman Brothers against the probability of at least one exceedance in the month ending two weeks before the bankruptcy. The right column plots show the corresponding posterior densities associated with the idiosyncratic component, and the middle column includes the posterior density associated with the systematic component. The vertical line indicates the mean of the posterior distribution.

the odds ratios of the probability of at least one exceedance in the month starting two weeks after the bankruptcy of Lehman Brothers versus the probability of at least one exceedance in the month ending two weeks before the bankruptcy. Note that all sectors show an increase in risk after the Lehman Brothers bankruptcy. However, the increase is lower for financials than it is for the other sectors (the estimated posterior probabilities are 1.000, 0.913 and 0.977 for consumer staples, energy and information technology, respectively). Indeed, note that systematic risk increases after the bankruptcy of Lehman Brothers but the idiosyncratic risk associated with financials actually decreases (as does the one for energy, although to a lesser degree), while the idiosyncratic risks associated with information technology and consumer staples increased. The increase in risk in the information technology and consumer staples sectors could be explained by the fact that one of the main effects of Lehman's bankruptcy was a collapse in the short-term corporate debt market. Hence, although the bankruptcy of Lehman Brothers actually reduced the uncertainty in the financial sector of the economy, it caused real damage to companies in other sectors that are extremely dependent on short-term debt. Note that companies that are part of the S&P500 energy sector are typically not reliant in short-term funding, hence the limited impact of Lehman's bankruptcy in their idiosyncratic risk.

4.1. Model validation. The model was validated using two different approaches. First, an out-of-sample cross-validation exercise was conducted, with cross-validation datasets being constructed by randomly selecting 20% of the observations from each sector to be used as held-out data. The remaining 80% of the data was used to fit our nonparametric model and generate nominal 90% highest posterior density (HPD) intervals for new exceedances. The true coverage of these HPD intervals was then evaluated on the held-out data. Figure 7 presents examples of cross-validation samples and the corresponding densities for two different sectors.

We repeated the process described above for 10 different cross-validation datasets, with the results presented in Figure 8. As expected, there is variability in the coverage rates depending on the sector and the specific cross-validation dataset. However, the results suggest that for the most part the real coverage rates are in line with the nominal coverage, which suggest that the model does not under- or over-fit.

In addition to the cross-validation exercise described above, in-sample goodness of fit was investigated using quantile-quantile plots for the posterior distribution of inter-event times. More specifically, we use the time-rescaling theorem [see, e.g., Daley and Vere-Jones (2003)], according to which if $\{t_{j,k} : k = 1, \dots, n_j\}$ is a realization from a nonhomogeneous Poisson process with cumulative intensity $\Lambda_j(t)$, then the transformed point pattern $\{\Lambda_j(t_{j,k}) : k = 1, \dots, n_j\}$ is a realization from a homogeneous Poisson process with unit intensity. Hence, if the Poisson process assumption is correct, then we would expect the transformed inter-arrival times

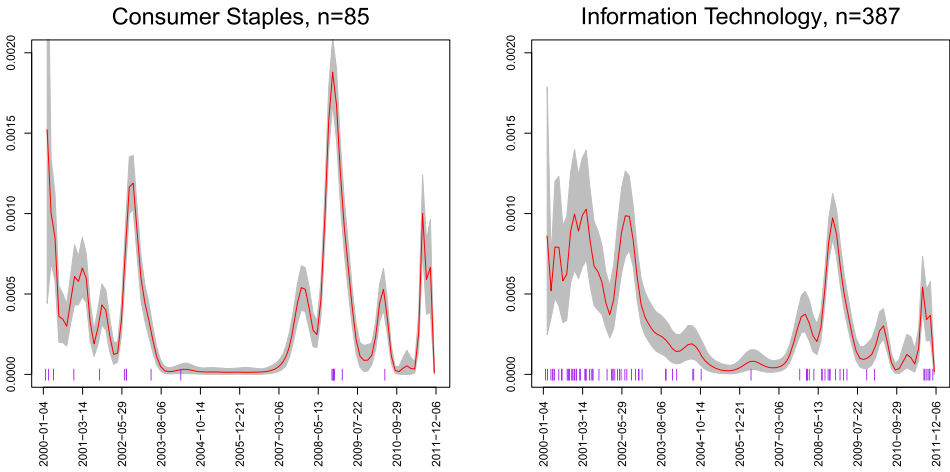


FIG. 7. Examples of cross-validation datasets (plotted on the horizontal axis) for two sectors, and the posterior mean and 95% pointwise interval estimates for the densities associated with them.

$z_{j,1}, \dots, z_{j,n_j}$ defined as $z_{j,k} = 1 - \exp\{-[\Lambda_j(t_{j,k}|G_j, \tau) - \Lambda_j(t_{j,k-1}|G_j, \tau)]\}$ (with the convention $t_{j,0} = 0$) to be uniformly distributed on the unit interval. A variant of this check for extreme value models was originally proposed by Hill (1975) in the context of threshold selection.

Figure 9 presents quantile-quantile plots of the posterior means of these transformed inter-arrival times, $E(z_{j,k}|\text{data})$, for $k = 1, \dots, n_j$, against the quantiles of

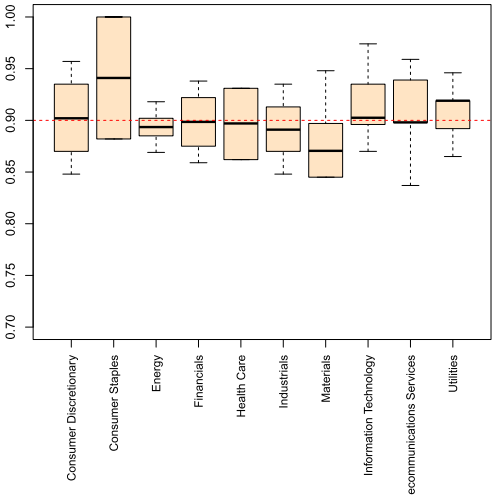


FIG. 8. Results from the cross-validation exercise to investigate the coverage rate of the highest posterior density intervals associated with the nonparametric model.

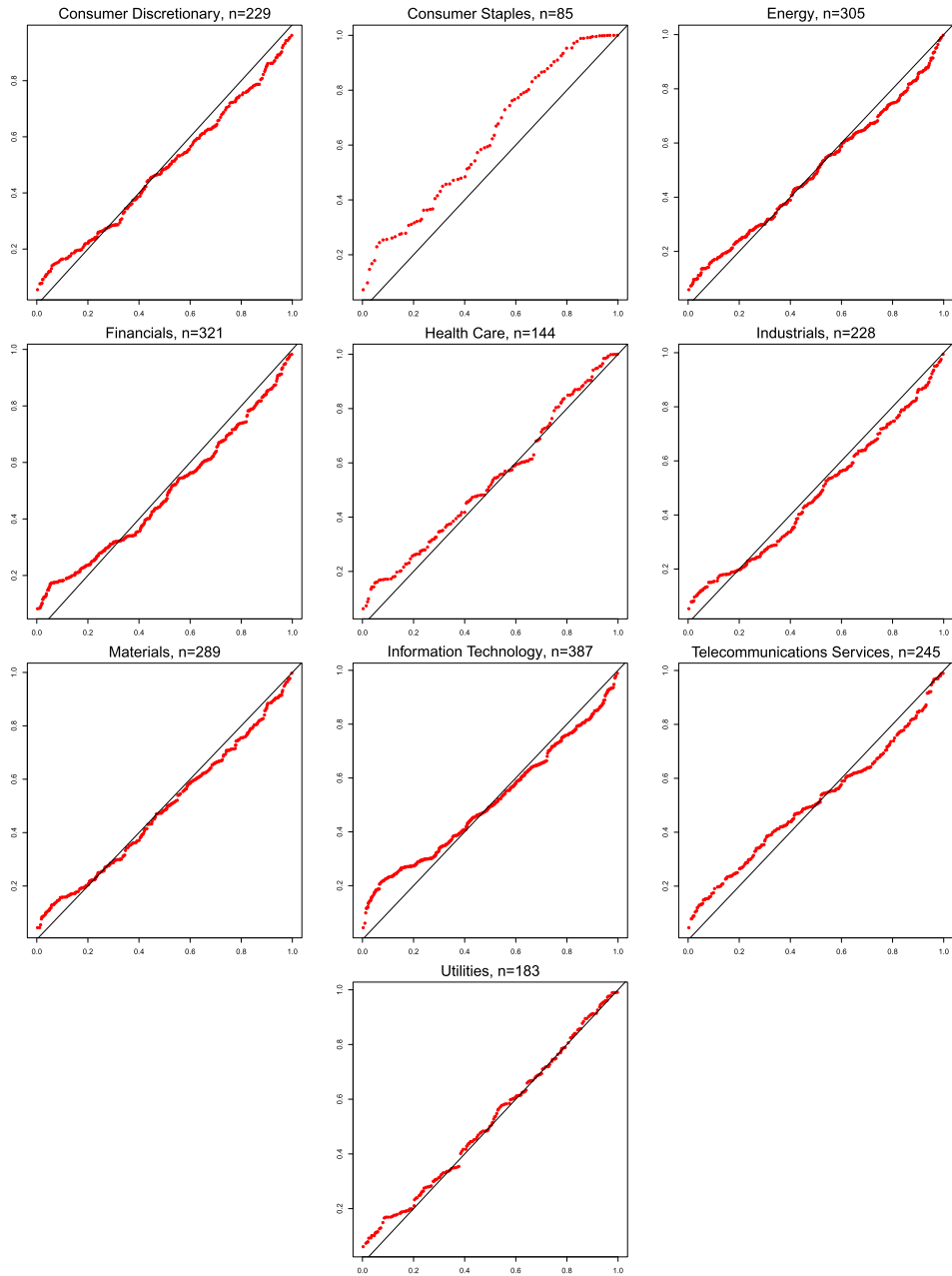


FIG. 9. *Quantile-quantile plot of the posterior means for the transformed inter-arrival times against the quantiles of a uniform distribution (solid line) for each of the ten S&P500 sectors.*

a uniform distribution for each of the ten S&P500 sectors. For the most part the fit appears acceptable, although there is some evidence of poor fit for a couple of sectors. In particular, note that for consumer staples our model tends to systematically predict shorter inter-arrival periods than those that would be expected under the Poisson model. Similar, albeit less dramatic, biases can also be seen for information technology. Note that these two sectors (along with telecommunications services) are identified by the cross-validation exercise as the ones where the model overfits the most.

4.2. Prior sensitivity analysis. In addition to testing the sensitivity of our analysis to the choice of the threshold u (see the supplementary material [Rodríguez, Wang and Kottas (2017)]), we carried out a comprehensive sensitivity analysis to assess the effect of prior distributions on posterior inferences. First, we considered three alternative sets of hyperparameters for the mixture kernel scale parameter τ corresponding to $\text{IGam}(5, 4000)$, $\text{IGam}(10, 7200)$ and $\text{IGam}(2, 500)$ priors. The hyperparameters were selected to represent a range of situations where the prior mean is both larger and smaller than the one used for our previous analysis, as well as different levels of concentration. Posterior inferences were mostly unaffected under any of these scenarios.

Next, we considered four alternative prior specifications for the precision parameters $\alpha_0, \dots, \alpha_{10}$, specifically, a $\text{Gam}(4, 1/3)$, a $\text{Gam}(10, 2)$, a $\text{Gam}(2, 0.4)$ and a $\text{Gam}(3, 3)$ prior. Inferences for the intensity function were mostly unchanged under these prior distributions. However, inferences for individual hyperparameters were somewhat affected. In particular, smaller values for $E(\alpha_j)$ naturally lead to somewhat smaller posterior means for the α_j , but also to larger posterior values for τ and an increase in the posterior mean for some of the ϵ_j [e.g., for consumer staples we have $\Pr(\epsilon_2 = 1 | \text{data}) = 0.71$ under the $\text{Gam}(3, 3)$ prior]. On the other hand, changes in the prior dispersion of α_j had no discernible effect on the individual posterior distributions of the hyperparameters, as long as the prior mean was kept constant.

To assess the effect of the baseline distribution H on posterior inferences, we considered an alternative rescaled beta distribution with $a_\mu = b_\mu = 3$. This prior tends to favor the localization of distress periods toward the middle of the time series. This alternative baseline distribution leads to somewhat smoother estimates for the density functions $f_1, \dots, f_{10}$, and to less multimodal estimates for the idiosyncratic intensities $f_1^*, \dots, f_{10}^*$. In addition, the posterior means for the α_j and for τ tend to be slightly smaller. However, the overall structure of the results is unchanged.

We also considered an alternative specification for the priors on γ_0^* and $\gamma_1^*, \dots, \gamma_{10}^*$ where we assume that the number of extreme returns in each sector is consistent with a 2% positive annualized return and a 25% annualized volatility for the S&P500, while only 50% of the exceedances come from the systematic component, and $\Pr(0.2 < \gamma_0^*(\gamma_0^* + \gamma_j^*)^{-1} < 0.8) = 0.99$. This leads to $a_{\gamma_0^*} = 7.65$,

$b_{\gamma_0^*} = 0.65$, and $a_{\gamma_j^*} = 7.65$, $b_{\gamma_j^*} = 0.65$ for $j \geq 1$. As before, these new priors have very little impact on the inference for the intensity functions. Also, although the posterior distributions on the weights $\epsilon_1, \dots, \epsilon_{10}$ are somewhat affected, the qualitative outcome is identical under both analyses. In particular, in both cases consumer staples is the only sector for which we find little or no evidence of an idiosyncratic component. The difference focuses on the level of evidence provided by the analysis: under the new prior, the idiosyncrasy test provides very strong evidence that shocks in consumer staples are driven exclusively by the systematic component [$\Pr(\epsilon_2 = 1|\text{data}) \approx 1$], while for the other sectors we have very strong evidence for the presence of idiosyncratic components [$\Pr(\epsilon_j = 1|\text{data}) = 0$ for $j = 1, 3, 4, \dots, 10$].

Finally, we investigated the effect on posterior inferences of alternative prior distributions on π . In addition to the original $\text{Beta}(0.5, 2)$ prior (which favors the hypothesis that most of the exceedances are generated by the systematic component of the model), we considered a uniform and a $\text{Beta}(0.5, 0.5)$ prior for π . While the $\text{Beta}(0.5, 0.5)$ had a negligible effect on posterior inferences, the use of a uniform prior led to an increase in the posterior mean for some of the ϵ_j [e.g., $\Pr(\epsilon_2 = 1|\text{data}) = 0.76$].

5. Discussion. We have developed a novel modeling approach for simultaneous risk assessment in multiple sectors of the U.S. economy that focuses on modeling the dependence in the appearance of large losses in the different sectors constituting the S&P500 index. Some of the advantages of the proposed methodology include its nonparametric nature, its focus on extreme returns rather than average returns, and the interpretability of the components of the model. These features lead to robust indexes and hypothesis tests for the presence of idiosyncratic risks in specific sectors of the economy.

Although our application used data from equity markets to understand the interdependence across different sectors of the economy, the proposed model also has potential applications in debt markets. In particular, reduced-form credit risk models [e.g., Lando (1998)] are also Cox process models for debt default. Therefore, our model could be extended to estimate the probability of default for firms within and across multiple economic sectors.

Acknowledgments. The authors wish to thank the Editor, Brendan Murphy and two reviewers for constructive feedback and for comments that improved the presentation of the material in the paper.

SUPPLEMENTARY MATERIAL

Supplement to “Assessing systematic risk in the S&P500 index between 2000 and 2011: A Bayesian nonparametric approach” (DOI: [10.1214/16-AOAS987SUPP](https://doi.org/10.1214/16-AOAS987SUPP); .pdf). The supplementary material contains results from sensitivity analysis to the choice of the threshold, as well as details for the MCMC

algorithm and on posterior inference. It is available as [Rodríguez, Wang and Kottas \(2017\)](#).

REFERENCES

- ACHARYA, V., ENGLE, R. and RICHARDSON, M. (2012). Capital shortfall: A new approach to ranking and regulating systemic risks. *Am. Econ. Rev.* **102** 59–64.
- ACHARYA, V. V., PEDERSEN, L. H., PHILIPPON, T. and RICHARDSON, M. P. (2010). Measuring systemic risk. Technical report, AFA 2011 Denver Meetings Paper.
- ADRIAN, T. and BRUNNERMEIER, M. (2010). Covar: A systemic risk contribution measure. Technical report, Princeton Univ., Princeton, NJ.
- ALLEN, D., GERRANS, P., SINGH, A. and POWELL, R. (2009). Quantile regression: Its application in investment analysis. *Journal of the Securities Institute of Australia* **1** 1–12.
- ANTONIAK, C. E. (1974). Mixtures of Dirichlet processes with applications to Bayesian nonparametric problems. *Ann. Statist.* **2** 1152–1174. [MR365969](#)
- BARNES, M. L. and HUGHES, W. A. (2002). A quantile regression analysis of the cross section of stock market returns. Technical report, Federal Reserve Bank of Boston.
- BASSETTI, F., CASARIN, R. and LEISEN, F. (2014). Beta-product dependent Pitman–Yor processes for Bayesian inference. *J. Econometrics* **180** 49–72. [MR3188911](#)
- CHANG, M. C., HUNG, J. C. and NIEH, C. C. (2011). Reexamination of capital asset pricing model (CAPM): An application of quantile regression. *African Journal of Business Management* **5** 12684–12690.
- COX, D. R. (1955). Some statistical methods connected with series of events. *J. R. Stat. Soc. Ser. B. Stat. Methodol.* **17** 129–157; discussion, 157–164. [MR0092301](#)
- DALEY, D. J. and VERE-JONES, D. (2003). *An Introduction to the Theory of Point Processes*, 2nd ed. Springer, New York. [MR1950431](#)
- DANIEL, K. D., HIRSLEIFER, D. and SUBRAHMANYAM, A. (2001). Overconfidence, arbitrage, and equilibrium asset pricing. *J. Finance* **56** 921–965.
- DIACONIS, P. and YLVISAKER, D. (1985). Quantifying prior opinion. In *Bayesian Statistics 2* (J. Bernardo, M. H. DeGroot, D. V. Lindley and A. F. M. Smith, eds.) 133–156. North-Holland, Amsterdam. [MR0862488](#)
- ESCOBAR, M. D. and WEST, M. (1995). Bayesian density estimation and inference using mixtures. *J. Amer. Statist. Assoc.* **90** 577–588. [MR1340510](#)
- FAMA, E. F. and FRENCH, K. R. (1992). The cross-section of expected stock returns. *J. Finance* **47** 427–465.
- FAMA, E. F. and FRENCH, K. R. (2004). The capital asset pricing model: Theory and evidence. *Journal of Economic Perspectives* **18** 25–46.
- FERGUSON, T. S. (1973). A Bayesian analysis of some nonparametric problems. *Ann. Statist.* **1** 209–230. [MR0350949](#)
- FRENCH, C. W. (2003). The Treynor capital asset pricing model. *Journal of Investment Management* **1** 60–72.
- GELFAND, A. E., KOTTAS, A. and MACEACHERN, S. N. (2005). Bayesian nonparametric spatial modeling with Dirichlet process mixing. *J. Amer. Statist. Assoc.* **100** 1021–1035. [MR2201028](#)
- GELMAN, A. and RUBIN, D. B. (1992). Inferences from iterative simulation using multiple sequences. *Statist. Sci.* **7** 457–472.
- GRIFFIN, J. E. and WALKER, S. G. (2011). Posterior simulation of normalized random measure mixtures. *J. Comput. Graph. Statist.* **20** 241–259.
- GRIFFITHS, R. C. and MILNE, R. K. (1978). A class of bivariate Poisson processes. *J. Multivariate Anal.* **8** 380–395. [MR0512608](#)

- HATJISPYROS, S. J., NICOLERIS, T. and WALKER, S. G. (2011). Dependent mixtures of Dirichlet processes. *Comput. Statist. Data Anal.* **55** 2011–2025. [MR2785111](#)
- HILL, B. M. (1975). A simple general approach to inference about the tail of a distribution. *Ann. Statist.* **3** 1163–1174. [MR0378204](#)
- ISHWARAN, H. and JAMES, L. F. (2001). Gibbs sampling methods for stick-breaking priors. *J. Amer. Statist. Assoc.* **96** 161–173. [MR1952729](#)
- ISHWARAN, H. and ZAREPOUR, M. (2000). Markov chain Monte Carlo in approximate Dirichlet and Beta two-parameter process hierarchical models. *Biometrika* **87** 371–390. [MR1782485](#)
- KALLI, M., GRIFFIN, J. E. and WALKER, S. G. (2011). Slice sampling mixture models. *Stat. Comput.* **21** 93–105. [MR2746606](#)
- KOLOSIATIS, M., GRIFFIN, J. E. and STEEL, M. F. J. (2013). On Bayesian nonparametric modelling of two correlated distributions. *Stat. Comput.* **23** 1–15. [MR3018346](#)
- KOTTAS, A. and SANSÓ, B. (2007). Bayesian mixture modeling for spatial Poisson process intensities, with applications to extreme value analysis. *J. Statist. Plann. Inference* **137** 3151–3163. [MR2365118](#)
- KOTTAS, A., WANG, Z. and RODRÍGUEZ, A. (2012). Spatial modeling for risk assessment of extreme values from environmental time series: A Bayesian nonparametric approach. *Environmetrics* **23** 649–662. [MR3019057](#)
- KOTTAS, A., BEHSETA, S., MOORMAN, D. E., POYNOR, V. and OLSON, C. R. (2012). Bayesian nonparametric analysis of neuronal intensity rates. *Journal of Neuroscience Methods* **203** 241–253.
- LANDO, D. (1998). On Cox processes and credit risk securities. *Review of Derivatives Research* **2** 99–120.
- LEISEN, F. and LIJOI, A. (2011). Vectors of two-parameter Poisson–Dirichlet processes. *J. Multivariate Anal.* **102** 482–495. [MR2755010](#)
- LEISEN, F., LIJOI, A. and SPANÓ, D. (2013). A vector of Dirichlet processes. *Electron. J. Stat.* **7** 62–90.
- LIJOI, A. and NIPOTI, B. (2014). A class of hazard rate mixtures for combining survival data from different experiments. *J. Amer. Statist. Assoc.* **109** 802–814. [MR3223751](#)
- LIJOI, A., NIPOTI, B. and PRÜNSTER, I. (2014a). Bayesian inference with dependent normalized completely random measures. *Bernoulli* **20** 1260–1291. [MR3217444](#)
- LIJOI, A., NIPOTI, B. and PRÜNSTER, I. (2014b). Dependent mixture models: Clustering and borrowing information. *Comput. Statist. Data Anal.* **71** 417–433. [MR3131980](#)
- LIJOI, A. and PRÜNSTER, I. (2010). Models beyond the Dirichlet process. In *Bayesian Nonparametrics* (N. L. Hjort, C. Holmes, P. Müller and S. G. Walker, eds.) 80–136. Cambridge Univ. Press, Cambridge.
- MACEACHERN, S. N. (2000). Dependent Dirichlet processes. Technical report, Dept. Statistics, Ohio State Univ., Columbus, OH.
- MANDELBROT, B. (1963). The variation of certain speculative prices. *Journal of Business* **36** 394–419.
- MÜLLER, P., QUINTANA, F. and ROSNER, G. (2004). A method for combining inference across related nonparametric Bayesian models. *J. R. Stat. Soc. Ser. B. Stat. Methodol.* **66** 735–749. [MR2088779](#)
- ROBERT, C. P. and CASELLA, G. (2004). *Monte Carlo Statistical Methods*, 2nd ed. Springer, New York. [MR2080278](#)
- RODRÍGUEZ, A., WANG, Z. and KOTTAS, A. (2017). Supplement to “Assessing systematic risk in the S& P500 index between 2000 and 2011: A Bayesian nonparametric approach.” DOI:[10.1214/16-AOAS987SUPP](#).
- SETHURAMAN, J. (1994). A constructive definition of Dirichlet priors. *Statist. Sinica* **4** 639–650.
- TADDY, M. A. (2010). Autoregressive mixture models for dynamic spatial Poisson processes: Application to tracking intensity of violent crime. *J. Amer. Statist. Assoc.* **105** 1403–1417. [MR2796559](#)

TADDY, M. A. and KOTTAS, A. (2012). Mixture modeling for marked Poisson processes. *Bayesian Anal.* **7** 335–362.

TREYNOR, J. L. (1961). Market Value, Time, and Risk. Unpublished manuscript.

TREYNOR, J. L. (1962). Toward a Theory of Market Value of Risky Assets. Unpublished manuscript.

XIAO, S., KOTTAS, A. and SANSÓ, B. (2015). Modeling for seasonal marked point processes: An analysis of evolving hurricane occurrences. *Ann. Appl. Stat.* **9** 353–382. [MR3341119](#)

A. RODRÍGUEZ

A. KOTTAS

DEPARTMENT OF APPLIED MATHEMATICS AND STATISTICS

UNIVERSITY OF CALIFORNIA, SANTA CRUZ

SANTA CRUZ, CALIFORNIA 95064

USA

E-MAIL: abel@soe.ucsc.edu

thanos@soe.ucsc.edu

Z. WANG

IAC PUBLISHING LABS

OAKLAND, CALIFORNIA 94607

USA

E-MAIL: ziweiwang520@gmail.com