



J. R. Statist. Soc. B (2019)
81, Part 4, pp. 781–804

Dynamic shrinkage processes

Daniel R. Kowal

Rice University, Houston, USA

and David S. Matteson and David Ruppert

Cornell University, Ithaca, USA

[Received July 2017. Final revision April 2019]

Summary. We propose a novel class of dynamic shrinkage processes for Bayesian time series and regression analysis. Building on a global–local framework of prior construction, in which continuous scale mixtures of Gaussian distributions are employed for both desirable shrinkage properties and computational tractability, we model dependence between the local scale parameters. The resulting processes inherit the desirable shrinkage behaviour of popular global–local priors, such as the horseshoe prior, but provide additional localized adaptivity, which is important for modelling time series data or regression functions with local features. We construct a computationally efficient Gibbs sampling algorithm based on a Pólya–gamma scale mixture representation of the process proposed. Using dynamic shrinkage processes, we develop a Bayesian trend filtering model that produces more accurate estimates and tighter posterior credible intervals than do competing methods, and we apply the model for irregular curve fitting of minute-by-minute Twitter central processor unit usage data. In addition, we develop an adaptive time varying parameter regression model to assess the efficacy of the Fama–French five-factor asset pricing model with momentum added as a sixth factor. Our dynamic analysis of manufacturing and healthcare industry data shows that, with the exception of the market risk, no other risk factors are significant except for brief periods.

Keywords: Asset pricing; Dynamic linear model; Stochastic volatility; Time series; Trend filtering

1. Introduction

The global–local class of prior distributions is a popular and successful mechanism for providing shrinkage and regularization in a broad variety of models and applications. Global–local priors use continuous scale mixtures of Gaussian distributions to produce desirable shrinkage properties, such as (approximate) sparsity or smoothness, often leading to highly competitive and computationally tractable estimation procedures. For example, in the variable-selection context, exact sparsity inducing priors such as the spike-and-slab prior become intractable for even a moderate number of predictors. By comparison, global–local priors that shrink towards sparsity, such as the horseshoe prior (Carvalho *et al.*, 2010), produce competitive estimators with greater scalability and are validated by theoretical results, simulation studies and a variety of applications (Datta and Ghosh, 2013; van der Pas *et al.*, 2014). Unlike non-Bayesian counterparts such as the lasso (Tibshirani, 1996), shrinkage priors also provide adequate uncertainty quantification for parameters of interest (Kyung *et al.*, 2010; van der Pas *et al.*, 2014).

Address for correspondence: Daniel R. Kowal, Department of Statistics, Rice University, 6100 Main Street, Houston, TX 77005-1892, USA.
E-mail: daniel.kowal@rice.edu

© 2019 The Authors Journal of the Royal Statistical Society: Series B 1369–7412/19/81781
(Statistical Methodology) Published by John Wiley & Sons Ltd on behalf of the Royal Statistical Society.
This is an open access article under the terms of the Creative Commons Attribution-NonCommercial License, which permits use, distribution and reproduction in any medium, provided the original work is properly cited and is not used for commercial purposes.

Global–local priors define a joint distribution for a set of variables $\{\omega_t\}_{t=1}^T$ and induce shrinkage via two key components: a *global* scale parameter τ , which controls the shrinkage that is common to all $\{\omega_t\}$, and *local* scale parameters $\{\lambda_t\}_{t=1}^T$, which control the shrinkage for each individual ω_t . Careful choice of priors for λ_t^2 and τ^2 provides both the flexibility to accommodate large signals and adequate shrinkage of noise (e.g. Carvalho *et al.* (2010)). Most commonly, the local scale parameters $\{\lambda_t\}$ are assumed to be *a priori* independent and identically distributed (IID). However, it can be advantageous to forgo the independence assumption. In the dynamic setting, in which the observations are time ordered, it is natural to allow the local scale parameter λ_t to depend on the history of the shrinkage process $\{\lambda_s\}_{s < t}$. Such model-based dependence may improve the ability of the model to adapt dynamically, which is important for time series estimation, forecasting and inference.

For a motivating example, consider the minute-by-minute Twitter central processor unit (CPU) usage data in Fig. 1 (James *et al.*, 2016). The data (Fig. 1(a)) show an overall smooth trend interrupted by irregular jumps throughout the morning and early afternoon, with an increase in volatility from 16.00 to 18.00 hours. It is important to identify both abrupt changes as well as slowly varying intraday trends. To model these features, consider the following Gaussian dynamic linear model (DLM):

$$\begin{aligned} y_t &= \beta_t + \epsilon_t, & [\epsilon_t | \sigma_t] &\stackrel{\text{indep}}{\sim} N(0, \sigma_t^2), \\ \Delta^2 \beta_{t+1} &= \omega_t, & [\omega_t | \tau, \lambda_t] &\stackrel{\text{indep}}{\sim} N(0, \tau^2 \lambda_t^2) \end{aligned} \quad (1)$$

where $\Delta^2 \beta_{t+1} = \Delta \beta_t - \Delta \beta_{t-1}$ for differencing operator Δ . Model (1) is a Bayesian adaptation of the *trend filtering* model of Kim *et al.* (2009) and Tibshirani (2014), also proposed by Faulkner and Minin (2018), and includes stochastic volatility (SV) for the observation error variance σ_t^2 with a global–local shrinkage prior for the second differences of the conditional mean, $\Delta^2 \beta_{t+1} = \omega_t$. The shrinkage behaviour of ω_t determines the path of β_t : when ω_t is pulled towards zero, β_t is locally linear, whereas large innovations $|\omega_t|$ correspond to large changes in the slope of β_t (Fig. 1(b)). Naturally, the global and local scale parameters τ and λ_t play a vital role in determining the magnitude of each ω_t , and thus the smoothness and adaptability of the conditional mean β_t .

Fig. 1 illustrates model (1) applied to the Twitter CPU usage data based on a *dynamic horseshoe prior* for $\{\omega_t\}$, which is introduced in Section 2. Unlike the classical horseshoe prior of Carvalho *et al.* (2010), the dynamic extension proposed incorporates dependence between the local scale parameters λ_t , so that the shrinkage behaviour at time t is informed by $\{\lambda_s\}_{s < t}$. Notably, the resulting posterior expectation of β_t and credible bands for the posterior predictive distribution of y_t adapt to both irregular jumps and smooth trends (Fig. 1(b)). The horseshoe-like shrinkage behaviour of λ_t is evident: values of λ_t are either near 0, corresponding to aggressive shrinkage of $\omega_t = \Delta^2 \beta_t$ to 0, or large, corresponding to large absolute changes in the slope of β_t (Fig. 1(d)). Importantly, Fig. 1 also provides motivation for a *dynamic* shrinkage process: there is clear *volatility clustering* of $\{\lambda_t\}$, in which the shrinkage that is induced by λ_t persists for consecutive time points. The volatility clustering reflects—and motivates—the temporally adaptive shrinkage behaviour of the dynamic shrinkage process.

More broadly, we introduce a framework for modelling dependence between the local scale parameters $\{\lambda_t\}$. Using a novel log-scale representation of a broad class of shrinkage priors, we gain the ability to incorporate a variety of widely successful models for dependent data, such as (vector) auto-regressions, linear regressions, Gaussian processes and factor models, in the shrinkage process model. Focusing on the case of dynamic dependence, we demonstrate that the

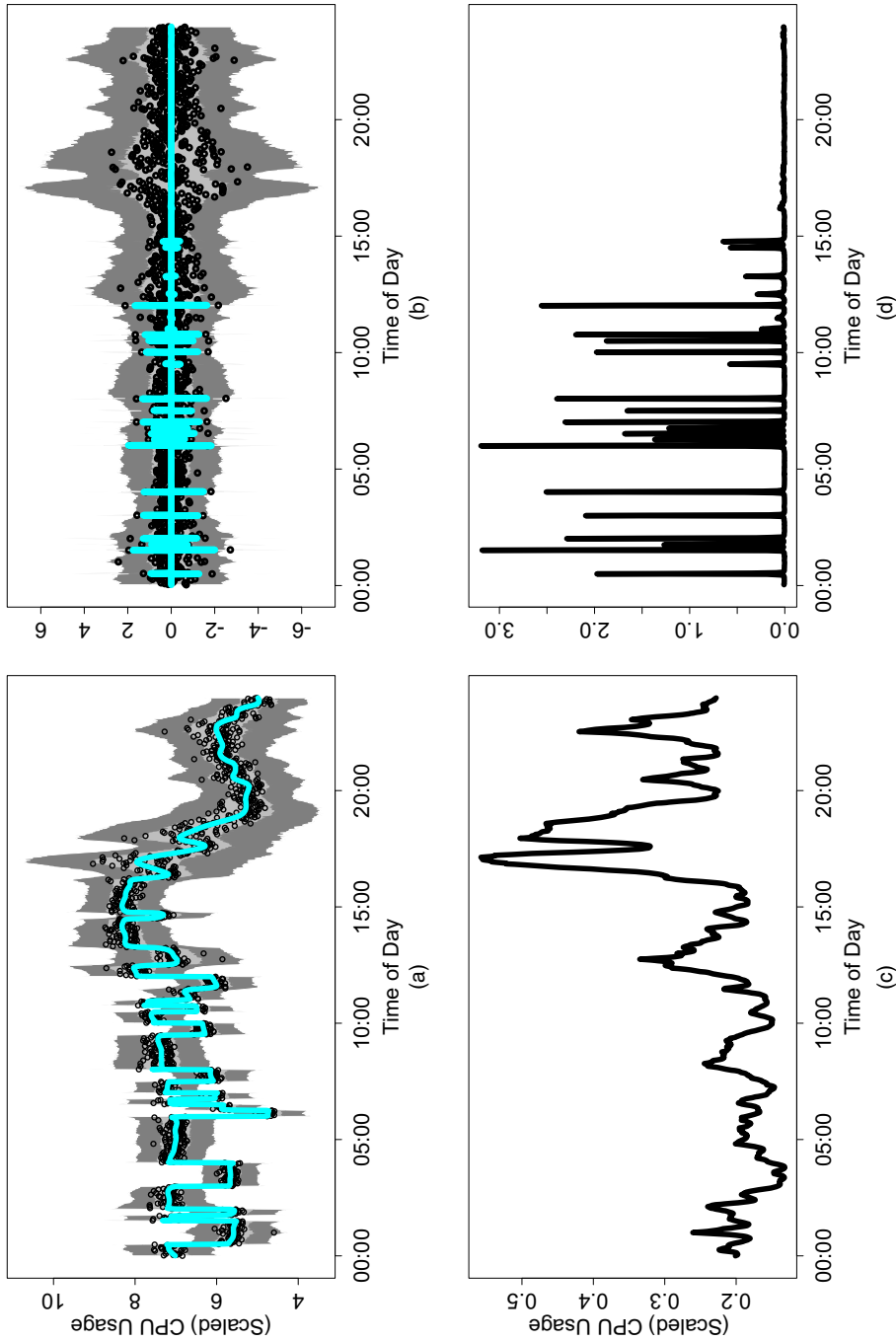


Fig. 1. Bayesian trend filtering with dynamic horseshoe process innovations of minute-by-minute Twitter CPU usage data: (a) observed data y_t (O), posterior expectation (—) of ω_t and 95% pointwise highest posterior density credible intervals (■) and 95% simultaneous credible bands (■) for the posterior predictive distribution of y_t ; (b) second difference of observed data, $\Delta^2 y_t$ (O), posterior expectation of $\omega_t = \Delta^2 \beta_t$ (—) and 95% pointwise highest posterior density credible intervals (■) and simultaneous credible bands (■) for the posterior predictive distribution of $\Delta^2 y_t$; (c) posterior expectation of time-dependent observation standard deviations, τ_t ; (d) posterior expectation of time-dependent innovation (prior) standard deviations, τ_t .

additional structure in the shrinkage process produces more accurate point estimates as well as substantially tighter and more adaptive credible intervals for both real and simulated data. To accompany the log-scale representation of global–local shrinkage priors proposed, we design new parameter expansion techniques using Pólya–gamma random variables for efficient and scalable posterior inference.

Shrinkage priors have been applied successfully and broadly for time series modelling. Belmonte *et al.* (2014), Korobilis (2013) and Bitto and Frühwirth-Schnatter (2019) proposed (global) shrinkage priors for DLMs and time series regression but did not allow for (local) time-specific shrinkage for each variable. Discrete mixture and spike-and-slab models offer an alternative approach but suffer from inherent computational challenges. One option is to restrict the space of models under consideration: Chan *et al.* (2012) included or excluded a variable for all times, whereas Frühwirth-Schnatter and Wagner (2010) also considered whether each variable is globally static or dynamic. The extensions in Huber *et al.* (2019) and Uribe and Lopes (2017) allowed for locally static or dynamic variables, whereas Nakajima and West (2013) provided a procedure for local thresholding of dynamic coefficients. Rockova and McAlinn (2017) developed an optimization approach for dynamic variable selection, which provides point estimates.

Perhaps most comparable with the methodology proposed, Kalli and Griffin (2014) proposed a class of priors which exhibit dynamic shrinkage by using normal–gamma auto-regressive processes. The Kalli and Griffin (2014) prior is a dynamic extension of the normal–gamma prior of Griffin and Brown (2010) and provides improvements in forecasting performance relative to non-dynamic shrinkage priors. However, the Kalli and Griffin (2014) model requires careful specification of several hyperparameters and hyperpriors, and the computation requires sophisticated adaptive Markov chain Monte Carlo (MCMC) techniques, which results in lengthy computation times. By comparison, our proposed class of dynamic shrinkage processes is far more general and includes the dynamic horseshoe process as a special case—which notably does not require tuning of hyperparameters. Empirically, for time varying parameter regression models with dynamic shrinkage, the Kalli and Griffin (2014) MCMC sampler requires several hours, whereas our proposed MCMC sampler runs in only a few minutes (see Section 4.1 for details and a comparison of these methods).

We introduce dynamic shrinkage processes in Section 2. In Section 3, we apply the processes to develop a more adaptive Bayesian trend filtering (BTF) model for irregular curve fitting. The procedure proposed is compared with state of the art alternatives through simulations and the CPU usage data. In Section 4, we propose a time varying parameter regression model with dynamic shrinkage processes for adaptive regularization and evaluate the model by using simulations and an asset pricing example. In Section 5, we discuss the details of the Gibbs sampling algorithm, and we conclude in Section 6. Proofs and additional details are in the on-line supplement.

2. Dynamic shrinkage processes

Model (1) features a global–local scale mixture of Gaussian distributions for the innovations:

$$[\omega_t | \tau, \lambda_t] \stackrel{\text{indep}}{\sim} N(0, \tau^2 \lambda_t^2). \quad (2)$$

Global–local priors are particularly well suited for sparse data: τ determines the global level of sparsity for $\{\omega_t\}_{t=1}^T$, whereas large λ_t allow large absolute deviations of ω_t from its prior mean (0) and small λ_t provide extreme shrinkage to 0. We propose to model dependence of the

Table 1. Special cases of the inverted beta prior

Parameters	Prior	Reference
$\alpha = \beta = \frac{1}{2}$	Horseshoe	Carvalho <i>et al.</i> (2010)
$\alpha = \frac{1}{2}, \beta = 1$	Strawderman–Berger	Strawderman (1971)
$\alpha = 1, \beta = c - 2, c > 0$	Normal–exponential–gamma	Griffin and Brown (2005)
$\alpha = \beta \rightarrow 0$	(Improper) normal–Jeffreys	Figueiredo (2003)

log-variance process $h_t = \log(\tau^2 \lambda_t^2)$ by using the general dependent data model

$$h_t = \mu + \psi_t + \eta_t, \quad \eta_t \stackrel{\text{IID}}{\sim} Z(\alpha, \beta, 0, 1), \quad (3)$$

where $\mu = \log(\tau^2)$ corresponds to the global scale parameter, $\psi_t + \eta_t = \log(\lambda_t^2)$ corresponds to the local scale parameter and $Z(\alpha, \beta, \mu_z, \sigma_z)$ denotes the *Z-distribution* with density function

$$[z] = \{\sigma_z B(\alpha, \beta)\}^{-1} \exp\{(z - \mu_z)/\sigma_z\}^\alpha [1 + \exp\{(z - \mu_z)/\sigma_z\}^{-(\alpha+\beta)}], \quad z \in \mathbb{R},$$

where $B(\cdot, \cdot)$ is the beta function. In model (3), the local scale parameter $\lambda_t = \exp\{(\psi_t + \eta_t)/2\}$ has two components: ψ_t , which models dependence (see below), and η_t , which corresponds to the usual IID (log-) local scale parameter. When $\psi_t = 0$, model (3) reduces to the static setting and implies an *inverted beta* prior for λ_t^2 (see Section 2.1). Notably, the class of priors that is represented in model (3) includes the important shrinkage distributions in Table 1, in each case extended to the dependent data setting.

The role of ψ_t in model (3) is to provide locally adaptive shrinkage by modelling dependence. For dynamic dependence, we propose the *dynamic shrinkage process*

$$h_{t+1} = \mu + \phi(h_t - \mu) + \eta_t, \quad \eta_t \stackrel{\text{IID}}{\sim} Z(\alpha, \beta, 0, 1), \quad (4)$$

which is equivalent to model (3) with $\psi_t = \phi(h_{t-1} - \mu)$. Relative to static shrinkage priors, model (4) adds only one parameter ϕ and reduces to the static setting when $\phi = 0$. Importantly, the proposed Gibbs sampler for the parameters in process (4) is linear in the number of time points, T , and therefore is scalable. Other examples of model (3) include linear regression, $\psi_t = \mathbf{z}_t' \alpha$ for a vector of predictors $\mathbf{z}_t$, Gaussian processes and various multivariate models (see expression (7) in Section 4). We focus on dynamic dependence, but our modelling framework and computational techniques may be extended to incorporate more general dependence between shrinkage parameters.

2.1. Log-scale representations of global–local priors

Models (3) and (4) do not automatically induce desirable locally adaptive shrinkage properties: we must consider appropriate distributions for μ and η_t . To illustrate this point, suppose that $\eta_t \stackrel{\text{IID}}{\sim} N(0, \sigma_\eta^2)$ in process (4), which is a common assumption in SV modelling (Kim *et al.*, 1998). For the likelihood $[y_t | \omega_t] \sim^{\text{indep}} N(\omega_t, 1)$ and the prior (2), the posterior expectation of ω_t is $\mathbb{E}[\omega_t | \{y_s\}, \tau] = (1 - \mathbb{E}[\kappa_t | \{y_s\}, \tau]) y_t$, where

$$\kappa_t \equiv \frac{1}{1 + \text{var}(\omega_t | \tau, \lambda_t)} = \frac{1}{1 + \tau^2 \lambda_t^2} \quad (5)$$

is the *shrinkage parameter*. As noted by Carvalho *et al.* (2010), $\mathbb{E}[\kappa_t | \{y_s\}, \tau]$ is interpretable as the amount of shrinkage towards 0 *a posteriori*: $\kappa_t \approx 0$ yields minimal shrinkage (for signals),

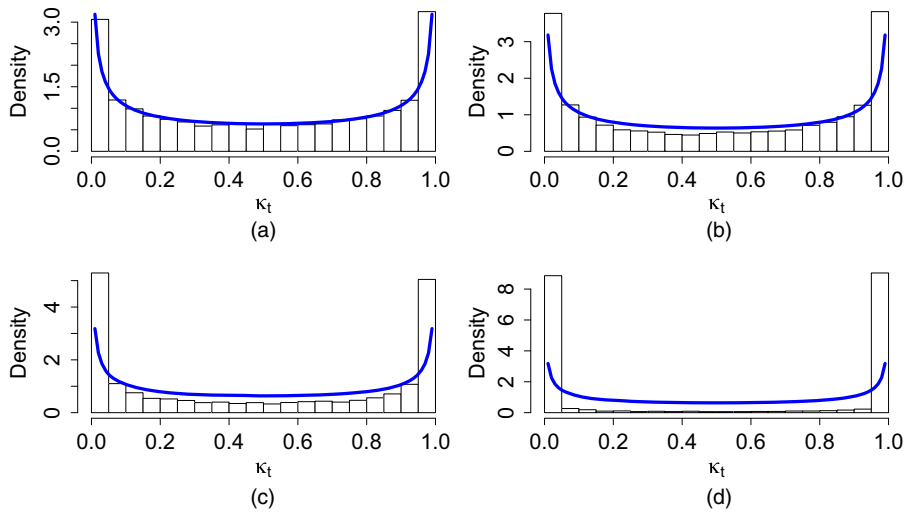


Fig. 2. Simulation-based estimate of the stationary distribution of κ_t for various AR(1) coefficients ϕ (—, density of κ_t in the static ($\phi = 0$) horseshoe, $[\kappa_t] \sim \text{beta}(\frac{1}{2}, \frac{1}{2})$): (a) $\phi = 0.25$; (b) $\phi = 0.5$; (c) $\phi = 0.75$; (d) $\phi = 0.99$

whereas $\kappa_t \approx 1$ yields maximal shrinkage to 0 (for noise). For the standard SV model and fixing $\phi = \mu = 0$ for simplicity, $\lambda_t = \exp(\eta_t/2)$ is log-normally distributed, and the shrinkage parameter has density

$$[\kappa_t] \propto \frac{1}{\kappa_t(1 - \kappa_t)} \exp \left\{ -\frac{1}{2\sigma_\eta^2} \log \left(\frac{1 - \kappa_t}{\kappa_t} \right)^2 \right\}.$$

Notably, the density for κ_t approaches 0 as $\kappa_t \rightarrow 0$ and as $\kappa_t \rightarrow 1$. As a result, direct application of the Gaussian SV model may overshrink true signals and undershrink noise.

By comparison, consider the horseshoe prior of Carvalho *et al.* (2010), which combines distribution (2) with $[\lambda_t] \sim \text{i.i.d. } C^+(0, 1)$, where C^+ denotes the half-Cauchy distribution. For fixed $\tau = 1$, the half-Cauchy prior on λ_t is equivalent to $\kappa_t \sim \text{i.i.d. } \text{beta}(\frac{1}{2}, \frac{1}{2})$, which induces a ‘horseshoe’ shape for the shrinkage parameter (Fig. 2). The horseshoe-like behaviour is ideal in sparse settings, since the prior density allocates most of its mass near 0 (minimal shrinkage of signals) and 1 (maximal shrinkage of noise).

To emulate the robustness and sparsity properties of the horseshoe and other shrinkage priors in the dynamic setting, we represent a general class of global–local shrinkage priors on the log-scale. As a motivating example, consider the special case of expressions (2) and (4) with $\phi = 0$, so $\log(\lambda_t^2) = \eta_t$. This example is illuminating: we equivalently express the (static) horseshoe prior by letting $\eta_t \stackrel{D}{=} \log(\lambda_t^2)$, where ‘ $\stackrel{D}{=}$ ’ denotes equality in distribution. In particular, $\lambda_t \sim C^+(0, 1)$ implies that $[\lambda_t^2] \propto (\lambda_t^2)^{-1/2} (1 + \lambda_t^2)^{-1}$ which implies that $[\eta_t] = \pi^{-1} \exp(\eta_t/2) \{1 + \exp(\eta_t)\}^{-1}$ is Z distributed with $\eta_t \sim Z(\frac{1}{2}, \frac{1}{2}, 0, 1)$. Importantly, Z -distributions may be written as mean–variance scale mixtures of Gaussian distributions (Barndorff-Nielsen *et al.*, 1982), which produces a useful framework for a parameter-expanded Gibbs sampler.

More generally, consider the inverted beta prior, denoted $\text{IB}(\beta, \alpha)$, for λ^2 with density $[\lambda^2] \propto (\lambda^2)^{\alpha-1} (1 + \lambda^2)^{-(\alpha+\beta)}$, $\lambda > 0$, using the parameterization of Polson and Scott (2012a, b). Special cases of the inverted beta distribution are provided in Table 1. This broad class of priors may be equivalently constructed via the variances λ_t^2 , the shrinkage parameters κ_t or the log-variances η_t .

Proposition 1. The following distributions are equivalent:

- (a) $\lambda^2 \sim \text{IB}(\beta, \alpha)$;
- (b) $\kappa = 1/(1 + \lambda^2) \sim \text{beta}(\beta, \alpha)$;
- (c) $\eta = \log(\lambda^2) = \log(\kappa^{-1} - 1) \sim Z(\alpha, \beta, 0, 1)$.

Note that the ordering of the parameters α and β is identical for the inverted beta and beta distributions but reversed for the Z -distribution.

Now consider the dynamic setting in which $\phi \neq 0$. Model (4) implies that the conditional prior variance for ω_t is

$$\exp(h_t) = \exp\{\mu + \phi(h_{t-1} - \mu) + \eta_t\} = \tau^2 \lambda_{t-1}^{2\phi} \tilde{\lambda}_t^2,$$

where $\tau^2 = \exp(\mu)$, $\lambda_{t-1}^2 = \exp(h_{t-1} - \mu)$ and $\tilde{\lambda}_t^2 = \exp(\eta_t) \sim \text{iID IB}(\beta, \alpha)$, as in the non-dynamic setting. This prior generalizes the $\text{IB}(\beta, \alpha)$ prior via the local variance term $\lambda_{t-1}^{2\phi}$, which incorporates information about the shrinkage behaviour at the previous time $t - 1$ in the prior for ω_t . We formalize the role of this local adjustment term with the following result.

Proposition 2. Suppose that $\eta \sim Z(\alpha, \beta, \mu_z, 1)$ for $\mu_z \in \mathbb{R}$. Then $\kappa = 1/\{1 + \exp(\eta)\} \sim \text{TPB}\{\beta, \alpha, \exp(\mu_z)\}$, where $\kappa \sim \text{TPB}(\beta, \alpha, \gamma)$ denotes the three-parameter beta (TPB) distribution with density $[\kappa] = B(\alpha, \beta)^{-1} \gamma^\beta \kappa^{\beta-1} (1 - \kappa)^{\alpha-1} \{1 + (\gamma - 1)\kappa\}^{-(\alpha+\beta)}$ for $\kappa \in (0, 1)$, $\gamma > 0$.

The TPB distribution (Armagan *et al.*, 2011) generalizes the beta distribution: $\gamma = 1$ produces the $\text{beta}(\beta, \alpha)$ distribution, whereas $\gamma > 1$ and $\gamma < 1$ allocate more mass near 0 and 1 respectively relative to the $\text{beta}(\beta, \alpha)$ distribution. In section B of the on-line supplementary material, we generalize for a Z -distribution with $\sigma_z \neq 1$, which produces a new class of shrinkage priors with additional flexibility in the distribution of κ , especially near 0 and 1.

For dynamic shrinkage processes, the TPB distribution arises as the conditional prior distribution of κ_{t+1} given $\{\kappa_s\}_{s \leq t}$.

Theorem 1. For the dynamic shrinkage process (4), the conditional prior distribution of the shrinkage parameter $\kappa_{t+1} = 1/(1 + \tau^2 \lambda_{t+1}^2)$ is $[\kappa_{t+1} | \{\kappa_s\}_{s \leq t}, \phi, \tau] \sim \text{TPB}[\beta, \alpha, \tau^{2(1-\phi)} \{(1 - \kappa_t)/\kappa_t\}^\phi]$ or, equivalently, $[\kappa_{t+1} | \{\lambda_s\}_{s \leq t}, \phi, \tau] \sim \text{TPB}(\beta, \alpha, \tau^2 \lambda_t^{2\phi})$.

The previous value of the shrinkage parameter κ_t , together with the AR(1) coefficient ϕ , informs the magnitude and direction of the distributional shift of κ_{t+1} .

Theorem 2. For the dynamic horseshoe process (4) with $\alpha = \beta = \frac{1}{2}$ and fixed $\tau = 1$, the conditional prior distribution in theorem 1 satisfies $\mathbb{P}(\kappa_{t+1} < \varepsilon | \{\kappa_s\}_{s \leq t}, \phi) \rightarrow 1$ as $\kappa_t \rightarrow 0$ for any $\varepsilon \in (0, 1)$ and fixed $\phi \neq 0$.

The mass of the conditional prior distribution for κ_{t+1} concentrates near 0—corresponding to minimal shrinkage of signals—when κ_t is near 0, so the shrinkage behaviour at time t informs the (prior) shrinkage behaviour at time $t + 1$.

We similarly characterize the posterior distribution of κ_{t+1} given $\{\kappa_s\}_{s \leq t}$ in the following theorem, which extends the results of Datta and Ghosh (2013) to the dynamic setting.

Theorem 3. Under the likelihood $[y_t | \omega_t] \sim \text{indep } N(\omega_t, 1)$, the prior (2) and the dynamic horseshoe process (4) with $\alpha = \beta = \frac{1}{2}$ and fixed $\phi \neq 0$, the posterior distribution of κ_{t+1} given $\{\kappa_s\}_{s \leq t}$ satisfies the following properties.

- (a) For any fixed $\varepsilon \in (0, 1)$, $\mathbb{P}(\kappa_{t+1} > 1 - \varepsilon | y_{t+1}, \{\kappa_s\}_{s \leq t}, \phi, \tau) \rightarrow 1$ as $\gamma_t \rightarrow 0$ uniformly in $y_{t+1} \in \mathbb{R}$, where $\gamma_t = \tau^{2(1-\phi)} \{(1 - \kappa_t)/\kappa_t\}^\phi$.

- (b) For any fixed $\varepsilon \in (0, 1)$ and $\gamma_t < 1$, $\mathbb{P}(\kappa_{t+1} < \varepsilon | y_{t+1}, \{\kappa_s\}_{s \leq t}, \phi, \tau) \rightarrow 1$ as $|y_{t+1}| \rightarrow \infty$.

Theorem 3, part (a), demonstrates that the posterior mass of $[\kappa_{t+1} | \{\kappa_s\}_{s \leq t}]$ concentrates near 1 as $\tau \rightarrow 0$, as in the non-dynamic horseshoe, but also as $\kappa_t \rightarrow 1$. Therefore, the dynamic horseshoe process provides an additional mechanism for shrinkage of noise, besides the global scale parameter τ , via the previous shrinkage parameter κ_t . Moreover, theorem 3, part (b), shows that, despite the additional shrinkage capabilities, the posterior mass of $[\kappa_{t+1} | \{\kappa_s\}_{s \leq t}]$ concentrates near 0 for large absolute signals $|y_{t+1}|$, which indicates responsiveness of the dynamic horseshoe process to large signals analogously to the static horseshoe prior.

When $|\phi| < 1$, the log-variance process $\{h_t\}$ is stationary, which implies that $\{\kappa_t\}$ is stationary. In Fig. 2, we plot a simulation-based estimate of the stationary distribution of κ_t for various values of ϕ under the dynamic horseshoe process. The stationary distribution of κ_t is similar to the static horseshoe distribution ($\phi = 0$) for $\phi < 0.5$, whereas for large values of ϕ the distribution becomes more peaked at 0 (less shrinkage of ω_t) and 1 (more shrinkage of ω_t). The result is intuitive: larger $|\phi|$ corresponds to greater persistence in shrinkage behaviour, so marginally we expect states of aggressive shrinkage or little shrinkage.

2.2. Scale mixtures via Pólya–gamma processes

For efficient computations, we develop a parameter expansion of model (3) by using a conditionally Gaussian representation for η_t . In doing so, we may incorporate Gaussian models—and accompanying sampling algorithms—for dependent data in model (3). Given a conditionally Gaussian parameter expansion, a Gibbs sampler for model (3) proceeds as follows:

- (a) draw the log-variances h_t , for which the conditional prior (3) is Gaussian, and
- (b) draw the parameters in μ and ψ_t , for which the conditional likelihood (3) is Gaussian.

For the log-variance sampler, we represent the likelihood for h_t on the log-scale and approximate the resulting distribution by using a known discrete mixture of Gaussian distributions (see Section 5). This approach is popular in SV modelling (e.g. Kim *et al.* (1998)), which is analogous to the dynamic shrinkage process (4). Importantly, the proposed parameter expansion inherits the computational complexity of the samplers for h_t and ψ_t : for the dynamic shrinkage processes (4), the proposed parameter expansion implies that the log-variance $\{h_t\}_{t=1}^T$ is a Gaussian DLM and therefore $\{h_t\}_{t=1}^T$ may be sampled jointly in $\mathcal{O}(T)$ computations (see Section 5).

The algorithm proposed requires a parameter expansion of $\eta_t \sim Z(\alpha, \beta, 0, 1)$ in process (4) as a scale mixture of Gaussian distributions. The representation of a Z-distribution as a mean–variance scale mixture of Gaussian distributions is due to Barndorff-Nielsen *et al.* (1982). For implementation, we build on the framework of Polson *et al.* (2013), who proposed a Pólya–gamma scale mixture of Gaussian distributions representation for Bayesian logistic regression. A Pólya–gamma random variable ξ with parameters $b > 0$ and $c \in \mathbb{R}$, denoted $\xi \sim \text{PG}(b, c)$, is an infinite convolution of gamma random variables:

$$\xi \stackrel{\text{D}}{=} (2\pi^2)^{-1} \sum_{k=1}^{\infty} g_k \left\{ \left(k - \frac{1}{2}\right)^2 - c^2 / (4\pi^2) \right\}^{-1}$$

where $g_k \sim \text{IID gamma}(b, 1)$. Properties of Pólya–gamma random variables may be found in Barndorff-Nielsen *et al.* (1982) and Polson *et al.* (2013). Our interest in Pólya–gamma random variables derives from their role in representing the Z-distribution as a mean–variance scale mixture of Gaussian distributions.

Theorem 4. The random variable $\eta \sim Z(\alpha, \beta, \mathbf{0}, \mathbf{1})$, or equivalently $\eta = \log(\lambda^2)$ with $\lambda^2 \sim \text{IB}(\beta, \alpha)$, is a mean–variance scale mixture of Gaussian distributions with

$$[\eta|\xi] \sim N\{\xi^{-1}(\alpha - \beta)/2, \xi^{-1}\}$$

and $[\xi] \sim \text{PG}(\alpha + \beta, 0)$. Moreover, the conditional distribution of ξ is $[\xi|\eta] \sim \text{PG}(\alpha + \beta, \eta)$.

When $\alpha = \beta$, the Z -distribution is symmetric and $\mathbb{E}[\eta|\xi] = 0$. Polson *et al.* (2013) proposed a sampling algorithm for Pólya–gamma random variables, which is extremely efficient when $b = 1$. In our setting, this corresponds to $\alpha + \beta = 1$, for which the horseshoe prior is the prime example. Importantly, this representation enables us to construct an efficient sampling algorithm that combines an $\mathcal{O}(T)$ sampling algorithm for the log-volatilities $\{h_t\}_{t=1}^T$ with a Pólya–gamma sampler for the mixing parameters.

3. Bayesian trend filtering with dynamic shrinkage processes

Dynamic shrinkage processes are particularly appropriate for DLMs. DLMs, such as model (1), combine an observation equation, which relates the observed data to latent state variables, and an evolution equation, which allows the state variables—and therefore the conditional mean of the data—to be dynamic. By construction, DLMs contain many parameters and therefore may benefit from structured regularization. The dynamic shrinkage processes proposed offer such regularization and, unlike existing methods, do so adaptively.

The DLM (1) may be generalized to a D th-order random-walk model: $\Delta^D \beta_{t+1} = \omega_t$, for $D \in \mathbb{Z}^+$ and $\beta_{t+1} = \omega_t \sim N(0, \tau^2 \lambda_t^2)$ for $t = 0, \dots, D-1$, where $D \in \mathbb{Z}^+$ is the degree of differencing. We refer to model (1) as a BTF model, with various choices available for the distribution of the innovation standard deviations $\tau \lambda_t$. Here, we propose a dynamic horseshoe process for the innovations ω_t . The aggressive shrinkage of the horseshoe prior forces small values of $|\omega_t| = |\Delta^D \beta_{t+1}|$ towards 0, whereas the responsiveness of the horseshoe prior permits large values of $|\Delta^D \beta_{t+1}|$. When $D = 2$, model (1) will shrink the conditional mean β_t towards a piecewise linear function with breakpoints determined adaptively, while allowing large absolute changes in the slopes. Further, using the *dynamic* horseshoe process, the shrinkage effects that are induced by λ_t are time dependent, which provides localized adaptability to regions with rapidly or slowly changing features. Following Carvalho *et al.* (2010) and Polson and Scott (2012b), we assume a half-Cauchy prior for the global scale parameter $\tau \sim C^+(0, \sigma_\epsilon/\sqrt{T})$, in which we scale by the observation error variance and the sample size (Piironen and Vehtari, 2016). Using Pólya–gamma mixtures, the implied conditional prior on $\mu = \log(\tau^2)$ is $[\mu|\sigma_\epsilon, \xi_\mu] \sim N\{\log(\sigma_\epsilon^2) - \log(T), \xi_\mu^{-1}\}$ with $\xi_\mu \sim \text{PG}(1, 0)$. We include the details of the Gibbs sampling algorithm for model (1) in Section 5, which is notably *linear* in the number of time points T : the full conditional posterior precision matrices for $\beta = (\beta_1, \dots, \beta_T)'$ and $\mathbf{h} = (h_1, \dots, h_T)'$ are D banded and tridiagonal respectively, which admit highly efficient $\mathcal{O}(T)$ back-band substitution sampling algorithms (see section C of the on-line supplement for empirical evidence).

3.1. Bayesian trend filtering: simulations

To assess the performance of the BTF model (1) with dynamic horseshoe innovations (model BTF-DHS), we compared the proposed methods with several competitive alternatives by using simulated data. We considered the following variations on BTF model (1): normal–inverse gamma innovations (model BTF-NIG) via $\tau^{-2} \sim \text{gamma}(0.001, 0.001)$ with $\lambda_t = 1$; (static) horseshoe priors for the innovations (model BTF-HS) via $\tau, \lambda_t \sim \text{i.i.d. } C^+(0, 1)$. In addition, we include the (non-Bayesian) trend filtering model of Tibshirani (2014) implemented by using

the R package `genlasso` (Arnold and Tibshirani, 2014), for which the regularization tuning parameter is chosen by using cross-validation (trend filtering). For all trend filtering models, we select $D = 2$, but the relative performance is similar for $D = 1$. Among non-trend filtering models, we include a smoothing spline estimator implemented via `smooth.spline()` in R (R Core Team, 2017); the wavelet-based estimator of Abramovich *et al.* (1998) (BayesThresh) implemented in the `wavethresh` package (Nason, 2016) and the nested Gaussian process model nGP of Zhu and Dunson (2013), which relies on a state space model framework for efficient computations, comparable with—but empirically less efficient than—the BTF model (1).

We simulated 100 data sets from the model $y_t = y_t^* + \epsilon_t$, where y_t^* is the true function and $\epsilon_t \sim_{\text{indep}} N(0, \sigma_*^2)$. We use the following true functions y_t^* from Donoho and Johnstone (1994): *Doppler*, *Bumps*, *Blocks* and *Heavisine*, implemented in the R package `wmtsa` (Constantine and Percival, 2016). The noise variance σ_*^2 is determined by selecting a root signal-to-noise ratio RSNR and computing $\sigma_* = \text{sd}(y_t^*)/\text{RSNR}$, where $\text{sd}(y_t^*)$ is the sample standard deviation of $\{y_t^*\}_{t=1}^T$. As in Zhu and Dunson (2013), we select $\text{RSNR} = 7$ and use a moderate length time series, $T = 128$.

In Fig. 3, we provide an example of each true curve y_t^* , together with the proposed BTF-DHS posterior expectations and credible bands. Notably, the BTF model (1) with dynamic horseshoe innovations provides an exceptionally accurate fit to each data set. Importantly, the posterior expectations and the posterior credible bands adapt to both slowly and rapidly changing behaviour in the underlying curves. The implementation is also efficient: the computation time for 10000 iterations of the Gibbs sampling algorithm, implemented in R (on a MacBook Pro, 2.7-GHz Intel Core i5), is about 45 s.

To compare the aforementioned procedures, we compute the root-mean-squared errors

$$\text{RMSE}(\hat{y}) = \sqrt{\left\{ \frac{1}{T} \sum_{t=1}^T (y_t^* - \hat{y}_t)^2 \right\}}$$

for all estimators $\hat{y}$ of the true function, y^* . The results are displayed in Fig. 4. The proposed BTF-DHS implementation substantially outperforms all competitors, especially for rapidly changing curves (*Doppler* and *Bumps*). The exceptional performance of BTF-DHS is paired with comparably small variability of RMSE, especially relative to the non-dynamic horseshoe model BTF-HS. Interestingly, the magnitude and variability of the RMSEs for BTF-DHS are related to the auto-regressive AR(1) coefficient ϕ : the 95% highest posterior density intervals (corresponding to Fig. 3) are (0.77, 0.97) (*Doppler*), (0.81, 0.97) (*Bumps*), (0.76, 0.96) (*Blocks*) and (−0.04, 0.74) (*Heavisine*). For the smoothest function, *Heavisine*, there is less separation between the estimators. Nonetheless, BTF-DHS performs the best, even though the highest posterior density interval for ϕ is wider and contains zero.

We are also interested in uncertainty quantification, and in particular how the dynamic horseshoe prior compares with the horseshoe prior. We compute the mean credible intervals widths

$$\text{MCIW} = \frac{1}{T} \sum_{t=1}^T (\hat{\beta}_t^{(97.5)} - \hat{\beta}_t^{(2.5)})$$

where $\hat{\beta}_t^{(97.5)}$ and $\hat{\beta}_t^{(2.5)}$ are the 97.5% and 2.5% quantiles respectively of the posterior distribution of β_t in model (1) for BTF-DHS and BTF-HS. The results are in Fig. 5. The dynamic horseshoe provides massive reductions in MCIW, again in all cases except for *Heavisine*, for which the methods perform similarly. Therefore, in addition to more accurate point estimation (Fig. 4), the BTF-DHS model produces significantly tighter credible intervals—while maintaining the correct nominal (frequentist) coverage.

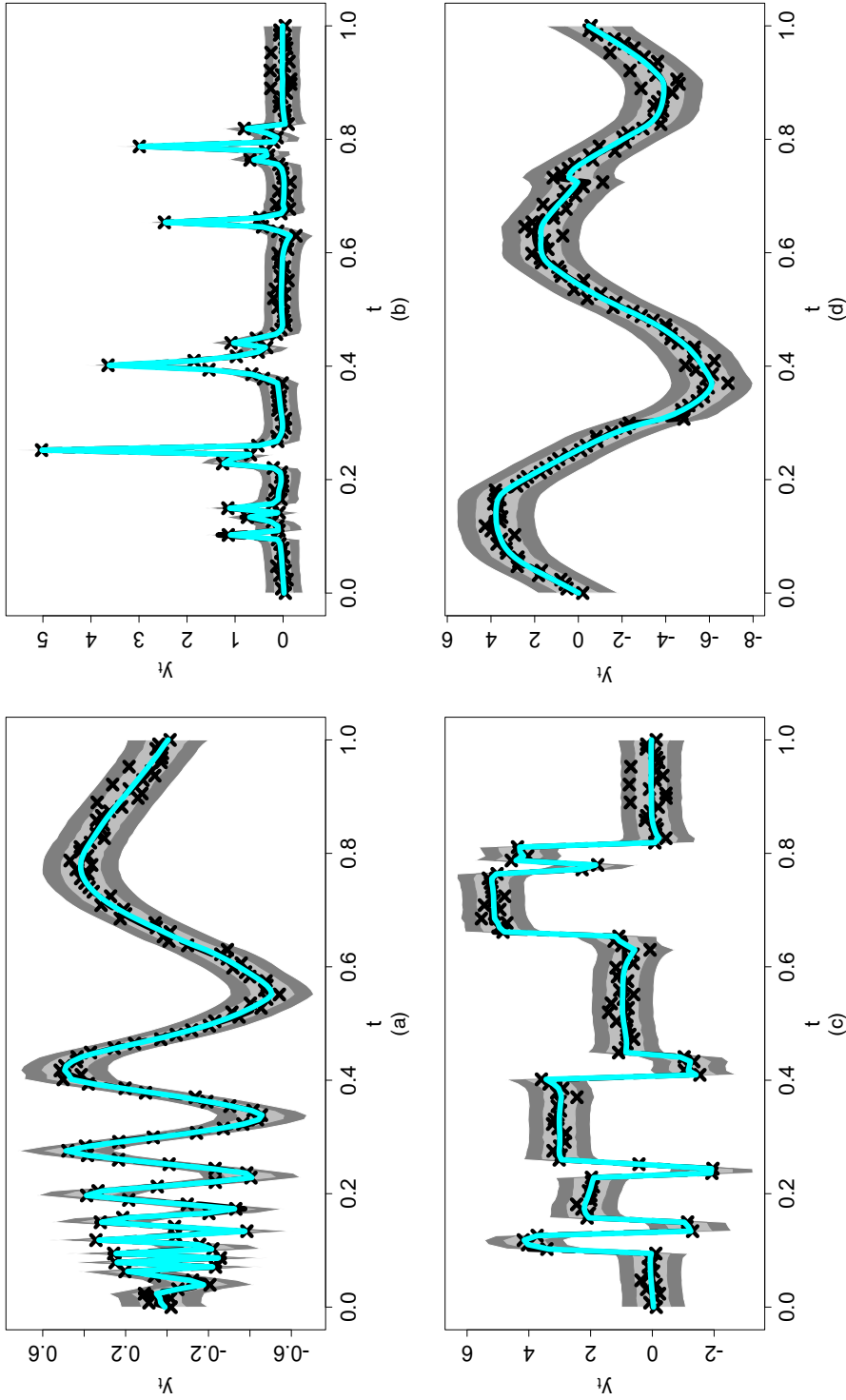


Fig. 3. Fitted curves for simulated data with $T = 128$ and $\text{RSNR} = 7$ (each panel includes the simulated observations (X), the posterior expectations of β_t (—) and the 95% pointwise highest posterior density credible intervals (■) and 95% simultaneous credible bands (■) for the posterior predictive distribution of $\{y_t\}$ under BTF-DHS model (1); the estimator proposed, as well as the uncertainty bands, accurately captures both slowly and rapidly changing behaviour in the underlying functions): (a) Doppler; (b) Bumps; (c) Blocks; (d) Heavisine

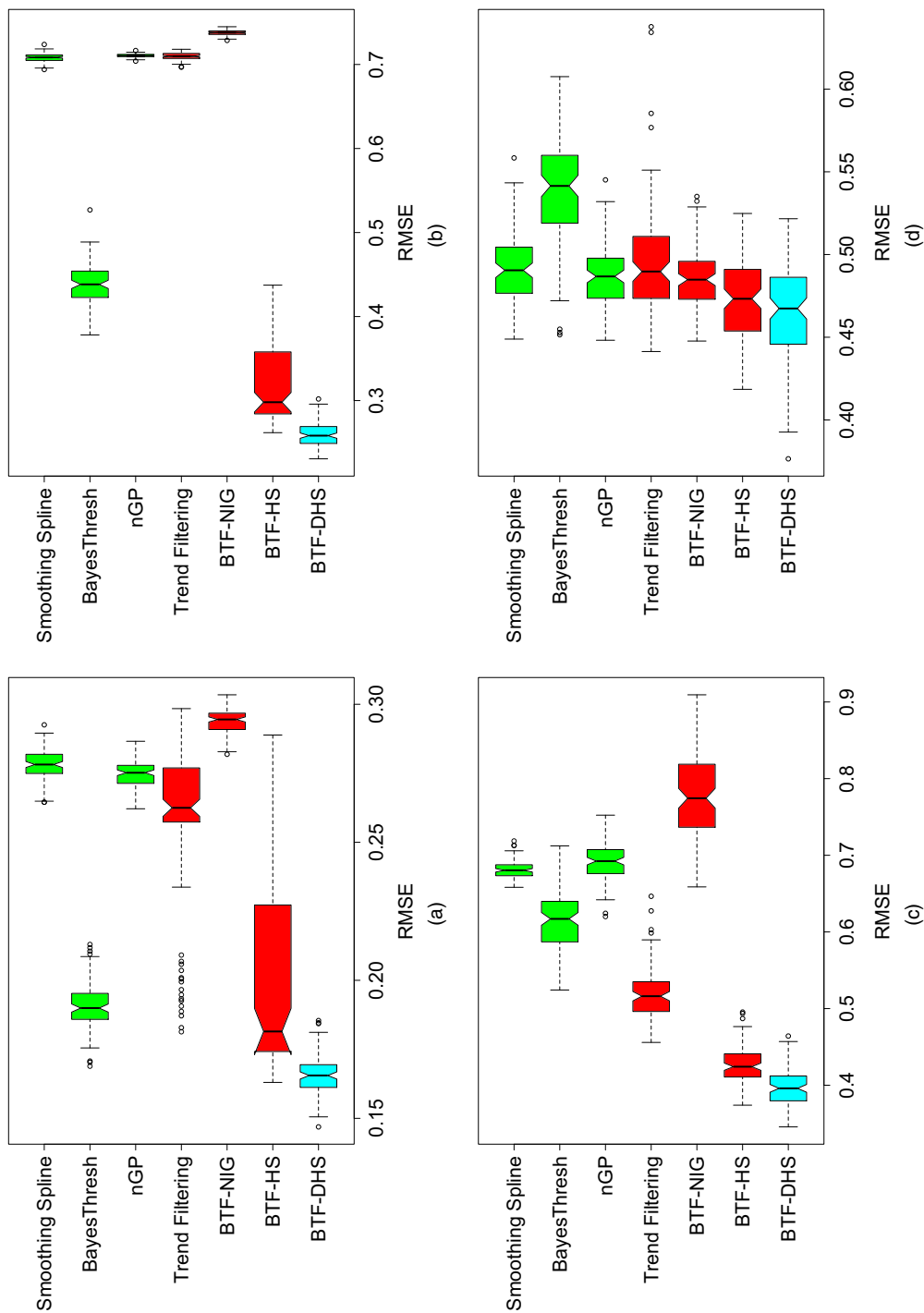


Fig. 4. Root-mean-squared errors for simulated data with $T = 128$ and $\text{RSNR} = 7$ (non-overlapping notches indicate significant differences between medians; the BTF estimators differ in their innovation distributions, which determine the shrinkage behaviour of the second-order differences; normal-inverse gamma, NIG; horseshoe, HS, and dynamic horseshoe, DHS: (a) Doppler; (b) Bumps; (c) Blocks; (d) Heavisine)

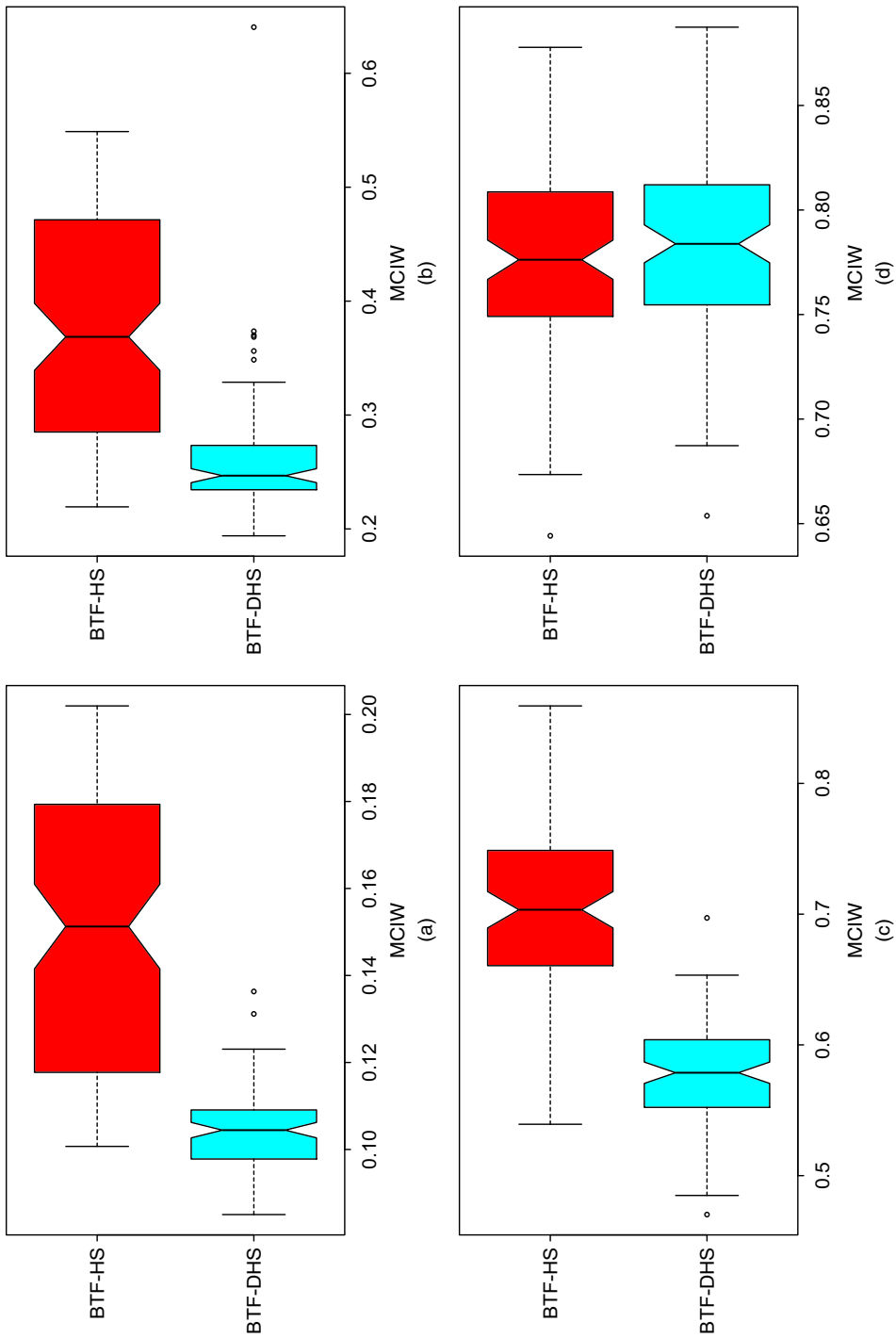


Fig. 5. Mean credible interval widths for simulated data with $T = 128$ and $\text{RSNR} = 7$ (non-overlapping notches indicate significant differences between medians; the BTF estimators differ in their innovation distributions, which determine the shrinkage behaviour of the second-order differences; normal-inverse gamma, NIG, horseshoe, HS, and dynamic horseshoe, DHS; (a) Doppler; (b) Bumps; (c) Blocks; (d) Heavisine)

3.2. Bayesian trend filtering: application to central processor unit usage data

To demonstrate the adaptability of the dynamic horseshoe process for model (1), we consider the CPU usage data in Fig. 1. The data exhibit substantial complexity: an overall smooth intraday trend but with multiple irregularly spaced jumps, and an increase in volatility from 16.00 to 18.00 hours. Our goal is to provide an accurate measure of the trend, including jumps, with appropriate uncertainty quantification. For this, we employ the BTF-DHS model (1), with a Gaussian AR(1) model on $\log(\sigma_t^2)$. For the additional sampling of the SV parameters, we use the algorithm of Kastner and Frühwirth-Schnatter (2014) implemented in the R package *stochvol* (Kastner, 2016).

We augment the simulation study of Section 3.1 with an out-of-sample comparison for the CPU usage data. We fit each model using 90% ($T = 1296$) of the data selected randomly for training and the remaining 10% ($T = 144$) for testing, which was repeated 100 times. Models were compared by using RMSE and MCIW.

Unlike the simulation study in Section 3.1, the subsampled data are *not* equally spaced. We employ a model-based imputation scheme for the unequally spaced data y_{t_i} , $i = 1, \dots, T$, which is similar to Elerian *et al.* (2001) and valid for missing observations. We expand the operative data set to include missing observations along an equally spaced grid, $t^* = 1, \dots, T^*$, such that, for each observation point i , $y_{t_i} = y_{t^*}$ for some t^* . Although $T^* \geq T$, possibly with $T^* \gg T$, all computations within the sampling algorithm, including the imputation sampling scheme for $\{y_{t^*} : t^* \neq t_i\}$, are linear in the number of (equally spaced) time points, T^* . Therefore, we may apply the same Gibbs sampling algorithm as before, with the additional step of drawing $y_{t^*} \sim^{\text{indep}} N(\beta_{t^*}, \sigma_{t^*}^2)$ for each unobserved $t^* \neq t_i$. Implicitly, this procedure assumes that the unobserved points are missing at random, which is satisfied by the aforementioned subsampling scheme.

The results of the out-of-sample estimation study are displayed in Fig. 6. The BTF procedures are notably superior to the non-BTF and smoothing spline estimators, and, as with the simulations of Section 3.1, the proposed BTF-DHS model substantially outperforms all competitors. Importantly, the significant reduction in MCIW by BTF-DHS indicates that the posterior credible intervals for the out-of-sample points y_{t^*} are substantially tighter for our method. By reducing uncertainty—while maintaining the approximately correct nominal (frequentist) coverage—the proposed BTF-DHS model provides greater power to detect local features. In addition, the MCMC algorithm for BTF-DHS is fast, despite the imputation procedure: 10000 iterations runs in about 80 s (in R on a MacBook Pro, 2.7-GHz Intel Core i5).

4. Joint shrinkage for time varying parameter models

Dynamic shrinkage processes are appropriate for multivariate time series and functional data models that may benefit from locally adaptive shrinkage properties. As outlined in Dangl and Halling (2012), models with *time varying parameters* are particularly important in financial and economic applications, because of the inherent structural changes in regulations, monetary policy, market sentiments and macroeconomic interrelationships that occur over time. Consider the following time varying parameter regression model with multiple dynamic predictors $\mathbf{x}_t = (x_{1,t}, \dots, x_{p,t})'$:

$$\begin{aligned} y_t &= \mathbf{x}_t' \boldsymbol{\beta}_t + \epsilon_t, & [\epsilon_t | \sigma_\epsilon] &\stackrel{\text{indep}}{\sim} N(0, \sigma_\epsilon^2), \\ \Delta^D \boldsymbol{\beta}_{t+1} &= \boldsymbol{\omega}_t, & [\boldsymbol{\omega}_{j,t} | \tau_0, \{\tau_k\}, \{\lambda_{k,s}\}] &\stackrel{\text{indep}}{\sim} N(0, \tau_0^2 \tau_j^2 \lambda_{j,t}^2) \end{aligned} \quad (6)$$

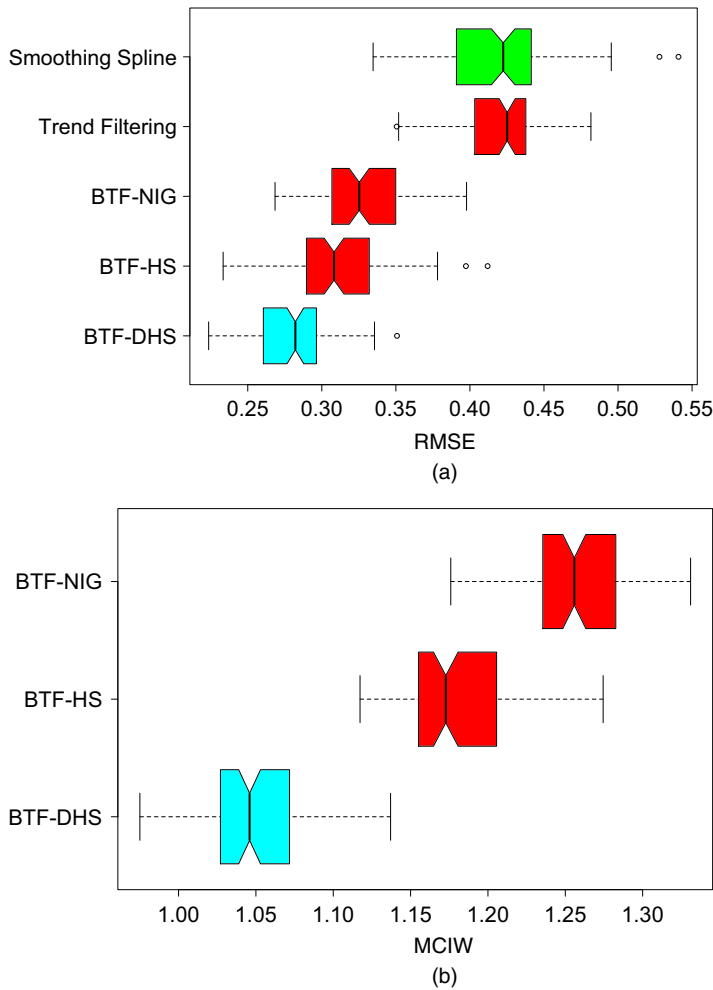


Fig. 6. (a) Root-mean-squared error and (b) mean credible interval widths for out-of-sample CPU usage data (non-overlapping notches indicate significant differences between medians; the BTF estimators differ in their innovation distributions, which determine the shrinkage behaviour of the second-order differences; normal-inverse gamma, NIG, horseshoe, HS, and dynamic horseshoe, DHS)

where $\beta_t = (\beta_{1,t}, \dots, \beta_{p,t})'$ is the vector of dynamic regression coefficients and $D \in \mathbb{Z}^+$ is the degree of differencing. Model (6) is also a (discretized) *concurrent functional linear model* (e.g. Ramsay and Silverman (2005)) and a *varying-coefficient model* (Hastie and Tibshirani, 1993) in the index t and therefore is broadly applicable. The prior for the innovations $\omega_{j,t}$ incorporates three levels of global–local shrinkage: a global shrinkage parameter τ_0 , a predictor-specific shrinkage parameter τ_j and a predictor- and time-specific local shrinkage parameter $\lambda_{j,t}$. Relative to existing time varying parameter regression models, our approach incorporates an additional layer of dynamic dependence: not only are the parameters time varying, but also the *relative influence* of the parameters is time varying via the shrinkage parameters—which are dynamically dependent themselves.

To provide jointly localized shrinkage of the dynamic regression coefficients $\{\beta_{j,t}\}$ analogous to the BTF model of Section 3, we extend process (4) to allow for multivariate time dependence

via a vector auto-regression on the log-variance:

$$\left. \begin{aligned} [\omega_{j,t} | \tau_0, \{\tau_k\}, \{\lambda_{k,s}\}] &\stackrel{\text{indep}}{\sim} N(0, \tau_0^2 \tau_j^2 \lambda_{j,t}^2), \\ h_{j,t} &= \log(\tau_0^2 \tau_j^2 \lambda_{j,t}^2), \quad j=1, \dots, p, \quad t=1, \dots, T, \\ \mathbf{h}_{t+1} &= \boldsymbol{\mu} + \boldsymbol{\Phi}(\mathbf{h}_t - \boldsymbol{\mu}) + \boldsymbol{\eta}_t, \quad \eta_{j,t} \stackrel{\text{iid}}{\sim} Z(\alpha, \beta, 0, 1), \end{aligned} \right\} \quad (7)$$

where $\mathbf{h}_t = (h_{1,t}, \dots, h_{p,t})'$, $\boldsymbol{\mu} = (\mu_1, \dots, \mu_p)'$, $\boldsymbol{\eta}_t = (\eta_{1,t}, \dots, \eta_{p,t})'$ and $\boldsymbol{\Phi}$ is the $p \times p$ vector auto-regression coefficient matrix. We assume that $\boldsymbol{\Phi} = \text{diag}(\phi_1, \dots, \phi_p)$ for simplicity, but non-diagonal extensions are available. Contemporaneous dependence may be introduced via a copula model for $\boldsymbol{\eta}_t$ but may reduce computational and MCMC efficiency. As in the univariate setting, we use Pólya–gamma mixtures for the log-variance evolution errors, $[\eta_{j,t} | \xi_{j,t}] \stackrel{\text{indep}}{\sim} N\{\xi_{j,t}^{-1}(\alpha - \beta)/2, \xi_{j,t}^{-1}\}$ with $\xi_{j,t} \sim \text{iid PG}(\alpha + \beta, 0)$ and $\alpha = \beta = \frac{1}{2}$. We augment model (7) with half-Cauchy priors for the predictor-specific and global parameters, $\tau_j \stackrel{\text{indep}}{\sim} C^+(0, 1)$ and $\tau_0 \sim C^+\{0, \sigma_\epsilon/\sqrt{(Tp)}\}$.

4.1. Time varying parameter models: simulations

We conducted a simulation study to evaluate competing variations of the time varying parameter regression model (6), in particular relative to the proposed dynamic shrinkage process DHS in expression (7). Similarly to the simulations of Section 3.1, we focus on the distribution of the innovations $\omega_{j,t}$ and again include the normal–inverse gamma (model NIG) and the (static) horseshoe HS as competitors, in each case selecting $D = 1$. We also include Belmonte *et al.* (2014), which uses the Bayesian lasso (model BL) as a prior on the innovations. Lastly, we include Kalli and Griffin (2014), which offers an alternative approach for dynamic shrinkage (model KG). Among models with non-dynamic regression coefficients, we include a lasso regression (Tibshirani, 1996) and an ordinary linear regression. These non-dynamic methods were non-competitive and have been excluded from the figures.

We simulated 100 data sets of length $T = 200$ from the model $y_t = \mathbf{x}_t' \boldsymbol{\beta}_t^* + \epsilon_t$, where the $p = 20$ predictors are $x_{1,t} = 1$ and $x_{j,t} \sim N(0, 1)$ for $j > 2$, and $\epsilon_t \sim N(0, \sigma_*^2)$ independently. We also consider auto-correlated predictors $x_{j,t}$ in section D.2 of the on-line supplement with similar results. The true regression coefficients $\boldsymbol{\beta}_t^* = (\beta_{1,t}^*, \dots, \beta_{p,t}^*)'$ are as follows: $\beta_{1,t}^* = 2$ is constant; $\beta_{2,t}^*$ is piecewise constant with $\beta_{2,t}^* = 0$ everywhere except $\beta_{2,t}^* = 2$ for $t = 41, \dots, 80$ and $\beta_{2,t}^* = -2$ for $t = 121, \dots, 160$; $\beta_{3,t}^* = (1/\sqrt{100}) \sum_{s=1}^t Z_s$ with $Z_s \sim \text{iid } N(0, 1)$ is a scaled random walk for $t \leq 100$ and $\beta_{3,t}^* = 0$ for $t > 100$; $\beta_{j,t}^* = 0$ for $j = 4, \dots, p = 20$. The predictor set contains a variety of functions: a constant non-zero function, a locally constant function, a slowly varying function that thresholds to 0 for $t > 100$, and 17 true 0s. The noise variance σ_*^2 is determined by selecting a root signal-to-noise ratio RSNR and computing $\sigma_* = \text{sd}(y_t^*)/\text{RSNR}$, where $y_t^* = \mathbf{x}_t' \boldsymbol{\beta}_t^*$ and $\text{sd}(y_t^*)$ is the sample standard deviation of $\{y_t^*\}_{t=1}^T$. We select RSNR = 3.

We evaluate competing methods using RMSEs for both y_t^* and $\boldsymbol{\beta}_t^*$ defined by

$$\text{RMSE}(\hat{y}) = \sqrt{\left\{ \frac{1}{T} \sum_{t=1}^T (y_t^* - \hat{y}_t)^2 \right\}}$$

and

$$\text{RMSE}(\hat{\beta}) = \sqrt{\left\{ \frac{1}{Tp} \sum_{t=1}^T \sum_{j=1}^p (\beta_{j,t}^* - \hat{\beta}_{j,t})^2 \right\}}$$

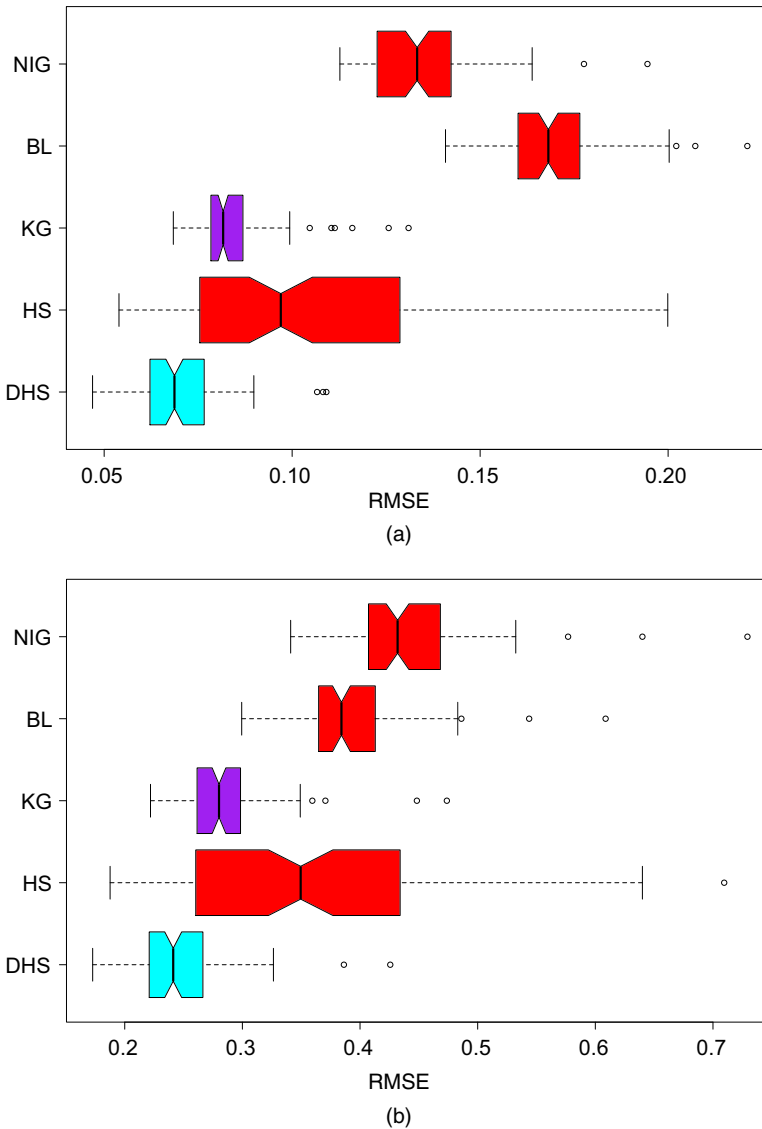


Fig. 7. (a) Root-mean-squared errors for the regression coefficients $\beta_{j,t}^*$ and (b) true curves, $y_t^* = \mathbf{x}_t' \beta_t^*$ for simulated data: non-overlapping notches indicate significant differences between medians

for all estimators $\hat{\beta}_t$ of the true regression functions, β_t^* with $\hat{y}_t = \mathbf{x}_t' \hat{\beta}_t$. The results are displayed in Fig. 7. The proposed BTF-DHS model substantially outperforms the competitors in both recovery of the true regression functions $\beta_{j,t}^*$ and estimation of the true curves y_t^* . Our closest competitor is Kalli and Griffin (2014), which also uses dynamic shrinkage, yet is less accurate in estimating the regression coefficients $\beta_{j,t}^*$ and the fitted values y_t^* . In addition, our MCMC algorithm is vastly more efficient: for 10000 MCMC iterations, the Kalli and Griffin (2014) sampler ran for 3 h 40 min (using MATLAB code from Professor Griffin's web site), whereas our algorithm completed in 6 min (on a MacBook Pro, 2.7-GHz Intel Core i5).

4.2. Time varying parameter models: the Fama–French asset pricing model

Asset pricing models commonly feature highly structured factor models to model the co-movement of stock returns parsimoniously. Such fundamental factor models identify common risk factors among assets, which may be treated as exogenous predictors in a time series regression. Popular approaches include the one-factor capital asset pricing model (Sharpe, 1964) and the three-factor Fama–French model FF-3 (Fama and French, 1993). Recently, the five-factor Fama–French model FF-5 (Fama and French, 2015) was proposed as an extension of FF-3 to incorporate additional common risk factors. However, outstanding questions remain regarding which, and how many, factors are necessary. Importantly, an attempt to address these questions must consider the dynamic component: the relevance of individual factors may change over time, particularly for different assets.

We apply model (6) to extend these fundamental factor models to the dynamic setting, in which the factor loadings are permitted to vary—perhaps rapidly—over time. For further generality, we append the momentum factor of Carhart (1997) to FF-5 to produce a fundamental factor model with six factors and dynamic factor loadings. Importantly, the shrinkage towards sparsity that is induced by the dynamic horseshoe process enables the model effectively to select out unimportant factors, which also may change over time. As in Section 3.2, we modify model (6) to include SV for the observation error.

To study various market sectors, we use weekly industry portfolio data from the web site of Kenneth R. French, which provide the value-weighted return of stocks in the given industry. We focus on manufacturing, Manuf, and healthcare, Hlth. For a given industry portfolio, the response variable is the returns in excess of the risk-free rate, $y_t = R_t - R_{F,t}$, with predictors $\mathbf{x}_t = (1, R_{M,t} - R_{F,t}, \text{SMB}_t, \text{HML}_t, \text{RMW}_t, \text{CMA}_t, \text{MOM}_t)'$, defined as follows: the *market risk factor* $R_{M,t} - R_{F,t}$ is the return on the market portfolio $R_{M,t}$ in excess of the risk-free rate $R_{F,t}$; the *size factor* SMB_t (small minus big) is the difference in returns between portfolios of small and large market value stocks; the *value factor* HML_t (high minus low) is the difference in returns between portfolios of high and low book-to-market value stocks; the *profitability factor* RMW_t is the difference in returns between portfolios of robust and weak profitability stocks; the *investment factor* CMA_t is the difference in returns between portfolios of stocks of low and high investment firms; the *momentum factor* MOM_t is the difference in returns between portfolios of stocks with high and low prior returns. These data are publicly available on Kenneth R. French's web site, which provides additional details on the portfolios. We standardize all predictors and the response to have unit variance.

We conduct inference on the time varying regression coefficients $\beta_{j,t}$ by using simultaneous credible bands. Unlike pointwise credible intervals, simultaneous credible bands control for multiple testing and may be computed as in Ruppert *et al.* (2003). Letting $B_{j,t}(\alpha)$ denote the $(1 - \alpha)\%$ simultaneous credible band for predictor j at time t , we compute the *simultaneous band score* (SIMBAS) (Meyer *et al.*, 2015), $P_{j,t} = \min\{\alpha : 0 \notin B_{j,t}(\alpha)\}$. The SIMBAS $P_{j,t}$ indicates the minimum level for which the simultaneous bands do not include zero, while controlling for multiple testing, and therefore may be used to detect which predictors j are important at time t . Globally, we compute *global Bayesian p-values* (GBPVs) (Meyer *et al.*, 2015), $P_j = \min_t\{P_{j,t}\}$ for each predictor j , which indicate whether or not a predictor is important at *any* time t . SIMBAS and GBPVs have proven useful in functional regression models but also are suitable for time varying parameter regression models to identify important predictors. We validate the selection ability of SIMBAS in section D.1 of the on-line supplementary material for the simulated data of Section 4.1.

In Fig. 8, we plot the posterior expectation and credible bands for the time varying regression coefficients and observation error SV for the weekly manufacturing industry data, from April

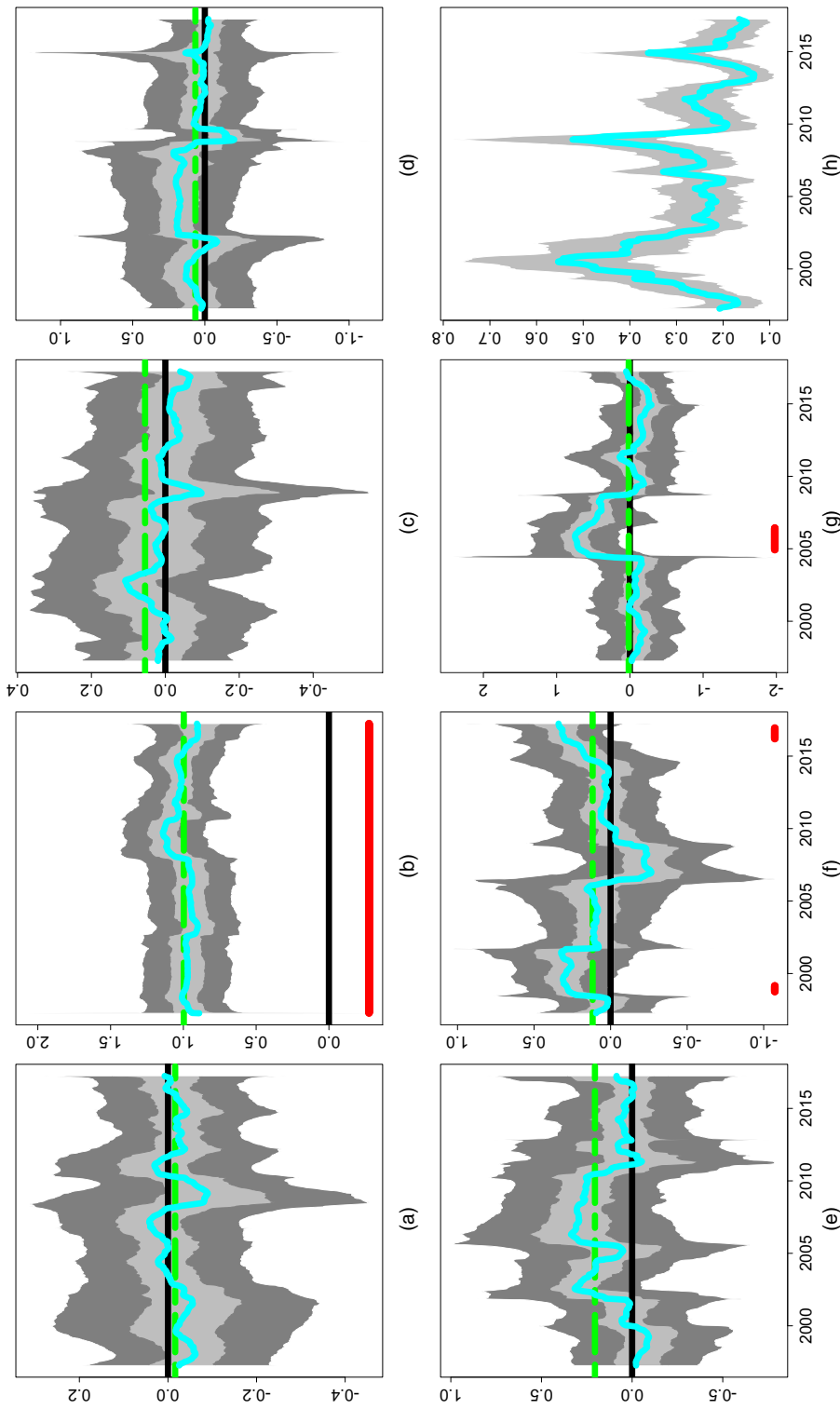


Fig. 8. Posterior expectations (—) for $\beta_{j,t}$ under the BTF-DHS model (6)–(7) for value-weighted manufacturing industry returns (—), ordinary linear regression estimate, and σ_t (—); 95% pointwise highest posterior density credible intervals (shaded gray) and 95% simultaneous credible bands (—) for $\beta_{j,t}$; (a) intercept; (b) Mkt-RF; (c) SMB; (d) HML; (e) RMW; (f) CMA; (g) MOM; (h) SV, σ_t

1st, 2007, to April 1st, 2017 ($T = 522$); an analogous plot for the healthcare industry data is in the on-line supplementary material (section E). For the manufacturing industry, the important factors that were identified by the GBPVs at the 5% level are the market risk ($R_{M,t} - R_{F,t}$; GBPV 0.000), investment (CMA_t ; GBPV 0.024) and momentum (MOM_t ; GBPV 0.019). However, the SIMBAS $P_{j,t}$ for CMA_t and MOM_t is below 5% only for brief periods (the red lines), which suggests that these important effects are intermittent. For the healthcare industry, the GBPVs identify market risk (GBPV 0.001) and value (HML_t ; GBPV 0.023) as the only important factors. Notably, the only common factor that is flagged by GBPVs in both the manufacturing and healthcare industries under model (6) over this time period is the market risk. This result suggests that the aggressive shrinkage behaviour of the dynamic shrinkage process is important in this setting, since several factors may be effectively irrelevant for some or all time points.

5. Markov chain Monte Carlo sampling algorithm and computational details

We design a Gibbs sampling algorithm for the dynamic shrinkage process. The sampling algorithm is both computationally and MCMC efficient, and builds on two main components:

- (a) a log-variance sampling algorithm (Kastner and Frühwirth-Schnatter, 2014) augmented with a Pólya–gamma sampler (Polson *et al.*, 2013);
- (b) a Cholesky factor algorithm (Rue, 2001) for sampling the state variables in a DLM.

Importantly, computations for each component are linear in the number of time points, which produces an efficient sampling algorithm.

The general sampling algorithm is as follows:

- (a) sample the dynamic shrinkage components (the log-volatilities $\{h_t\}$, the Pólya–gamma mixing parameters $\{\xi_t\}$, the unconditional mean of the log-variance μ , the AR(1) coefficient of the log-variance ϕ and the discrete mixture component indicators $\{s_t\}$);
- (b) sample the state variables $\{\beta_t\}$;
- (c) sample the observation error variance σ_ϵ^2 .

We provide details of the dynamic shrinkage process sampling algorithm in Section 5.1 and include the details for sampling steps (b) and (c) in the on-line supplement (section C).

5.1. Efficient sampling for the dynamic shrinkage process

Consider the (univariate) dynamic shrinkage process (4) with the Pólya–gamma parameter expansion of theorem 4. We provide implementation details for the dynamic horseshoe process with $\alpha = \beta = \frac{1}{2}$, but extensions to other cases are straightforward. The sampling framework of Kastner and Frühwirth-Schnatter (2014) represents the likelihood for h_t in prior (2) on the log-scale, and approximates the ensuing $\log(\chi_1^2)$ -distribution for the errors via a known discrete mixture of Gaussian distributions. In particular, let $\tilde{y}_t = \log(\omega_t^2 + c)$, where c is a small offset to avoid numerical issues. Conditionally on the mixture component indicators s_t , the likelihood is $\tilde{y}_t \sim^{\text{indep}} N(h_t + m_{s_t}, v_{s_t})$ where m_i and v_i , $i = 1, \dots, 10$, are the prespecified mean and variance components of the 10-component Gaussian mixture that was provided in Omori *et al.* (2007). Under model (4), the evolution equation is $h_{t+1} = \mu + \phi(h_t - \mu) + \eta_t$ with initialization $h_1 = \mu + \eta_0$ and innovations $[\eta_t | \xi_t] \sim^{\text{indep}} N(0, \xi_t^{-1})$ for $[\xi_t] \sim^{\text{IID}} \text{PG}(1, 0)$. Note that model (3) provides a more general setting, which similarly may be combined with the Gaussian likelihood for $\tilde{y}_t$ above.

To sample $\mathbf{h} = (h_1, \dots, h_T)$ jointly, we directly compute the posterior distribution of $\mathbf{h}$ and exploit the tridiagonal structure of the resulting posterior precision matrix. In particular, we equivalently have $\tilde{\mathbf{y}} \sim N(\mathbf{m} + \tilde{\mathbf{h}}, \Sigma_v)$ and $\mathbf{D}_\phi \mathbf{h} \sim N(\mathbf{0}, \Sigma_\xi)$, where $\mathbf{m} = (m_{s_1}, \dots, m_{s_T})'$, $\tilde{\mathbf{h}} = (h_1 - \mu, \dots, h_T - \mu)'$, $\tilde{\mu} = (\mu, (1 - \phi)\mu, \dots, (1 - \phi)\mu)'$, $\Sigma_v = \text{diag}(\{v_{s_t}\}_{t=1}^T)$, $\Sigma_\xi = \text{diag}(\{\xi_t^{-1}\}_{t=1}^T)$ and $\mathbf{D}_\phi$ is a lower triangular matrix with 1s on the diagonal, $-\phi$ on the first off-diagonal and 0s elsewhere. We sample from the posterior distribution of $\mathbf{h}$ by sampling from the posterior distribution of $\tilde{\mathbf{h}}$ and setting $\mathbf{h} = \tilde{\mathbf{h}} + \mu \mathbf{1}$ for $\mathbf{1}$ a T -dimensional vector of 1s. The required posterior distribution is $\tilde{\mathbf{h}} \sim N(\mathbf{Q}_h^{-1} \mathbf{l}_h, \mathbf{Q}_h^{-1})$, where $\mathbf{Q}_h = \Sigma_v^{-1} + \mathbf{D}_\phi' \Sigma_\xi^{-1} \mathbf{D}_\phi$ is a tridiagonal symmetric matrix with diagonal elements $\mathbf{d}_0(\mathbf{Q}_h)$ and first off-diagonal elements $\mathbf{d}_1(\mathbf{Q}_h)$ defined as

$$\mathbf{d}_0(\mathbf{Q}_h) = ((v_{s_1}^{-1} + \xi_1 + \phi^2 \xi_2), \dots, (v_{s_{T-1}}^{-1} + \xi_{T-1} + \phi^2 \xi_T), (v_{s_T}^{-1} + \xi_T)),$$

$$\mathbf{d}_1(\mathbf{Q}_h) = ((-\phi \xi_2), \dots, (-\phi \xi_{T-1}))$$

and

$$\mathbf{l}_h = \Sigma_v^{-1} (\tilde{\mathbf{y}} - \mathbf{m} - \tilde{\mu})$$

$$= \left(\frac{\tilde{y}_1 - m_{s_1} - \mu}{v_{s_1}}, \frac{\tilde{y}_2 - m_{s_2} - (1 - \phi)\mu}{v_{s_2}}, \dots, \frac{\tilde{y}_T - m_{s_T} - (1 - \phi)\mu}{v_{s_T}} \right)'.$$

Drawing from this posterior distribution is straightforward and efficient, using band back-substitution described in Kastner and Frühwirth-Schnatter (2014):

- compute the Cholesky decomposition $\mathbf{Q}_h = \mathbf{L}\mathbf{L}'$, where $\mathbf{L}$ is lower triangle;
- solve $\mathbf{L}\mathbf{a} = \mathbf{l}_h$ for $\mathbf{a}$;
- solve $\mathbf{L}'\tilde{\mathbf{h}} = \mathbf{a} + \mathbf{e}$ for $\tilde{\mathbf{h}}$, where $\mathbf{e} \sim N(\mathbf{0}, \mathbf{I}_T)$.

Conditionally on the log-volatilities $\{h_t\}$, we sample the AR(1) evolution parameters: the log-innovation precisions $\{\xi_t\}$, the auto-regressive coefficient ϕ and the unconditional mean μ . The precisions are distributed $[\xi_t | \eta_t] \sim \text{PG}(1, \eta_t)$ for $\eta_t = h_{t+1} - \mu - \phi(h_t - \mu)$, which we sample by using Polson *et al.* (2013). The Pólya-gamma sampler is efficient: using only exponential and inverse Gaussian draws, Polson *et al.* (2013) constructed an accept-reject sampler for which the probability of acceptance is uniformly bounded below at 0.99919, which does not require any tuning. Next, we assume the prior $(\phi + 1)/2 \sim \text{beta}(a_\phi, b_\phi)$, which restricts $|\phi| < 1$ for stationarity, and sample from the full conditional distribution of ϕ by using the slice sampler of Neal (2003). We select $a_\phi = 10$ and $b_\phi = 2$, which place most of the mass for the density of ϕ in $(0, 1)$ with a prior mean of $\frac{2}{3}$ and a prior mode of $\frac{4}{5}$ to reflect the likely presence of persistent volatility clustering. The prior for the global scale parameter is $\tau \sim C^+(0, \sigma_\epsilon/\sqrt{T})$, which implies that $\mu = \log(\tau^2)$ is $[\mu | \sigma_\epsilon, \xi_\mu] \sim N\{\log(\sigma_\epsilon^2/T), \xi_\mu^{-1}\}$ with $\xi_\mu \sim \text{PG}(1, 0)$. Including the initialization $h_1 \sim N(\mu, \xi_0^{-1})$ with $\xi_0 \sim \text{PG}(1, 0)$, the posterior distribution for μ is $\mu \sim N(\mathbf{Q}_\mu^{-1} \mathbf{l}_\mu, \mathbf{Q}_\mu^{-1})$ with $\mathbf{Q}_\mu = \xi_\mu + \xi_0 + (1 - \phi)^2 \sum_{t=1}^{T-1} \xi_t$ and $\mathbf{l}_\mu = \xi_\mu \log(\sigma_\epsilon^2/T) + \xi_0 h_1 + (1 - \phi) \sum_{t=1}^{T-1} \xi_t (h_{t+1} - \phi h_t)$. Sampling ξ_μ and ξ_0 follows the Pólya-gamma sampling scheme above.

Finally, we sample the discrete mixture component indicators s_t . The discrete mixture probabilities are straightforward to compute: the prior mixture probabilities are the mixing proportions that were given by Omori *et al.* (2007) and the likelihood is $\tilde{y}_t \sim^{\text{indep}} N(h_t + m_{s_t}, v_{s_t})$; see Kastner and Frühwirth-Schnatter (2014) for details.

6. Discussion and future work

Dynamic shrinkage processes provide a computationally convenient and widely applicable mechanism for incorporating adaptive shrinkage and regularization in existing models. By ex-

tending a broad class of global–local shrinkage priors to the dynamic setting, the resulting processes inherit the desirable shrinkage behaviour, but with greater time localization. The success of dynamic shrinkage processes suggests that other priors may benefit from log-scale or other appropriate representations, with or without additional dependence modelling.

As demonstrated in Sections 3 and 4, dynamic shrinkage processes are particularly appropriate for DLMs, including trend filtering and time varying parameter regression. In both settings, the DLMs with dynamic horseshoe innovations outperform all competitors in simulated data and produce reasonable and interpretable results for real data applications. Dynamic shrinkage processes may be useful in other DLMs, such as for modelling seasonality or change points. Given the exceptional curve fitting capabilities of the BTF model (1) with dynamic horseshoe innovations (BTF-DHS), a natural extension would be to incorporate BTF-DHS in more general additive, functional or longitudinal data models to capture irregular or local curve features. Similarly, models (1) and (6), as well as the dependent shrinkage of model (3), may be extended for multivariate responses to provide both contemporaneous and dynamic shrinkage.

Another promising area for applications of the methodology proposed is compressive sensing and signal processing, which commonly rely on approximations for estimation and prediction (e.g. Ziniel and Schniter (2013) and Wang *et al.* (2016)). The linear time complexity of our MCMC algorithm for BTF-DHS may offer the computational scalability to provide full Bayesian inference, and perhaps improved prediction and uncertainty quantification, which is notably absent from Ziniel and Schniter (2013) and Wang *et al.* (2016).

Acknowledgements

The authors thank the Joint Editor, Associate Editor and referees for very helpful comments. Financial support from National Science Foundation grant AST-1312903 (Kowal and Ruppert), a Xerox PARC Faculty Research award (Kowal and Matteson), National Science Foundation grant DMS-1455172, Cornell University Atkinson Center for a Sustainable Future grant AVF-2017, US Agency for International Development and the Cornell University Institute of Biotechnology and NYSTAR is gratefully acknowledged.

References

- Abramovich, F., Sapatinas, T. and Silverman, B. W. (1998) Wavelet thresholding via a Bayesian approach. *J. R. Statist. Soc. B*, **60**, 725–749.
- Armagan, A., Clyde, M. and Dunson, D. B. (2011) Generalized Beta mixtures of Gaussians. In *Advances in Neural Information Processing Systems*, pp. 523–531.
- Arnold, T. B. and Tibshirani, R. J. (2014) genlasso: path algorithm for generalized lasso problems. *R Package Version 1.3*.
- Barndorff-Nielsen, O., Kent, J. and Sørensen, M. (1982) Normal variance-mean mixtures and z distributions. *Int. Statist. Rev.*, **50**, 145–159.
- Belmonte, M. A., Koop, G. and Korobilis, D. (2014) Hierarchical shrinkage in time-varying parameter models. *J. Forecast.*, **33**, 80–94.
- Bitto, A. and Frühwirth-Schnatter, S. (2019) Achieving shrinkage in a time-varying parameter model framework. *J. Econometr.*, **210**, 75–97.
- Carhart, M. M. (1997) On persistence in mutual fund performance. *J. Finan.*, **52**, 57–82.
- Carvalho, C. M., Polson, N. G. and Scott, J. G. (2010) The horseshoe estimator for sparse signals. *Biometrika*, **97**, 465–480.
- Chan, J. C., Koop, G., Leon-Gonzalez, R. and Strachan, R. W. (2012) Time varying dimension models. *J. Bus. Econ. Statist.*, **30**, 358–367.
- Constantine, W. and Percival, D. (2016) wmtsa: wavelet methods for time series analysis. *R Package Version 2.0-2*.
- Dangl, T. and Halling, M. (2012) Predictive regressions with time-varying coefficients. *J. Finan. Econ.*, **106**, 157–181.
- Datta, J. and Ghosh, J. K. (2013) Asymptotic properties of Bayes risk for the horseshoe prior. *Bayes Anal.*, **8**, 111–132.

- Donoho, D. L. and Johnstone, J. M. (1994) Ideal spatial adaptation by wavelet shrinkage. *Biometrika*, **81**, 425–455.
- Elerian, O., Chib, S. and Shephard, N. (2001) Likelihood inference for discretely observed nonlinear diffusions. *Econometrica*, **69**, 959–993.
- Fama, E. F. and French, K. R. (1993) Common risk factors in the returns on stocks and bonds. *J. Finan. Econ.*, **33**, 3–56.
- Fama, E. F. and French, K. R. (2015) A five-factor asset pricing model. *J. Finan. Econ.*, **116**, 1–22.
- Faulkner, J. R. and Minin, V. N. (2018) Locally adaptive smoothing with Markov random fields and shrinkage priors. *Bayesn Anal.*, **13**, 225–252.
- Figueiredo, M. A. (2003) Adaptive sparseness for supervised learning. *IEEE Trans. Pattn Anal. Mach. Intell.*, **25**, 1150–1159.
- Frühwirth-Schnatter, S. and Wagner, H. (2010) Stochastic model specification search for Gaussian and partial non-Gaussian state space models. *J. Econometr.*, **154**, 85–100.
- Griffin, J. E. and Brown, P. J. (2005) Alternative prior distributions for variable selection with very many more variables than observations. *Technical Report*. Centre for Research in Statistical Methodology, University of Warwick, Coventry.
- Griffin, J. E. and Brown, P. J. (2010) Inference with normal-gamma prior distributions in regression problems. *Bayesn Anal.*, **5**, 171–188.
- Hastie, T. and Tibshirani, R. (1993) Varying-coefficient models (with discussion). *J. R. Statist. Soc. B*, **55**, 757–796.
- Huber, F., Kastner, G. and Feldkircher, M. (2019) Should I stay or should I go?: A latent threshold approach to large-scale mixture innovation models. *J. Appl. Econometr.*, to be published.
- James, N. A., Kejariwal, A. and Matteson, D. S. (2016) Leveraging cloud data to mitigate user experience from ‘Breaking Bad’. In *Proc. Int. Conf. Big Data* (eds J. Joshi, G. Karypis, L. Liu, X. Hu, R. Ak, Y. Xia, W. Xu, A.-H. Sato, S. Rachuri, L. Ungar, P. S. Yu, R. Govindaraju and T. Suzumura), pp. 3499–3508. New York: Institute of Electrical and Electronics Engineers.
- Kalli, M. and Griffin, J. E. (2014) Time-varying sparsity in dynamic regression models. *J. Econometr.*, **178**, 779–793.
- Kastner, G. (2016) Dealing with stochastic volatility in time series using the R package *stochvol*. *J. Statist. Softw.*, **69**, no. 5, 1–30.
- Kastner, G. and Frühwirth-Schnatter, S. (2014) Ancillarity-sufficiency interweaving strategy (ASIS) for boosting MCMC estimation of stochastic volatility models. *Computnl Statist. Data Anal.*, **76**, 408–423.
- Kim, S.-J., Koh, K., Boyd, S. and Gorinevsky, D. (2009) ℓ_1 trend filtering. *SIAM Rev.*, **51**, 339–360.
- Kim, S., Shephard, N. and Chib, S. (1998) Stochastic volatility: likelihood inference and comparison with ARCH models. *Rev. Econ. Stud.*, **65**, 361–393.
- Korobilis, D. (2013) Hierarchical shrinkage priors for dynamic regressions with many predictors. *Int. J. Forecast.*, **29**, 43–59.
- Kyung, M., Gill, J., Ghosh, M. and Casella, G. (2010) Penalized regression, standard errors, and Bayesian lassos. *Bayesn Anal.*, **5**, 369–411.
- Meyer, M. J., Coull, B. A., Versace, F., Cinciripini, P. and Morris, J. S. (2015) Bayesian function-on-function regression for multilevel functional data. *Biometrics*, **71**, 563–574.
- Nakajima, J. and West, M. (2013) Bayesian analysis of latent threshold dynamic models. *J. Bus. Econ. Statist.*, **31**, 151–164.
- Nason, G. (2016) *wavethresh: wavelets statistics and transforms*. *R Package Version 4.6.8*. Department of Mathematics, University of Bristol, Bristol.
- Neal, R. M. (2003) Slice sampling. *Ann. Statist.*, **31**, 705–741.
- Omori, Y., Chib, S., Shephard, N. and Nakajima, J. (2007) Stochastic volatility with leverage: fast and efficient likelihood inference. *J. Econometr.*, **140**, 425–449.
- van der Pas, S., Kleijn, B. and van der Vaart, A. (2014) The horseshoe estimator: posterior concentration around nearly black vectors. *Electron. J. Statist.*, **8**, 2585–2618.
- Piironen, J. and Vehtari, A. (2016) On the hyperprior choice for the global shrinkage parameter in the horseshoe prior. *Preprint arXiv:1610.05559*. Aalto University, Aalto.
- Polson, N. G. and Scott, J. G. (2012a) Local shrinkage rules, Lévy processes and regularized regression. *J. R. Statist. Soc. B*, **74**, 287–311.
- Polson, N. G. and Scott, J. G. (2012b) On the half-Cauchy prior for a global scale parameter. *Bayesn Anal.*, **7**, 887–902.
- Polson, N. G., Scott, J. G. and Windle, J. (2013) Bayesian inference for logistic models using Pólya–Gamma latent variables. *J. Am. Statist. Ass.*, **108**, 1339–1349.
- Ramsey, J. and Silverman, B. (2005) *Functional Data Analysis*. Berlin: Springer.
- R Core Team (2017) *R: a Language and Environment for Statistical Computing*. Vienna: R Foundation for Statistical Computing.
- Rockova, V. and McAlinn, K. (2017) Dynamic variable selection with spike-and-slab process priors. *Preprint arXiv:1708.00085*. Booth School of Business, University of Chicago, Chicago.
- Rue, H. (2001) Fast sampling of Gaussian Markov random fields. *J. R. Statist. Soc. B*, **63**, 325–338.
- Ruppert, D., Wand, M. P. and Carroll, R. J. (2003) *Semiparametric Regression*. New York: Cambridge University Press.

- Sharpe, W. F. (1964) Capital asset prices: a theory of market equilibrium under conditions of risk. *J. Finan.*, **19**, 425–442.
- Strawderman, W. E. (1971) Proper bayes minimax estimators of the multivariate normal mean. *Ann. Math. Statist.*, **42**, 385–388.
- Tibshirani, R. (1996) Regression shrinkage and selection via the lasso. *J. R. Statist. Soc. B*, **58**, 267–288.
- Tibshirani, R. J. (2014) Adaptive piecewise polynomial estimation via trend filtering. *Ann. Statist.*, **42**, 285–323.
- Uribe, P. V. and Lopes, H. F. (2017) Dynamic sparsity on dynamic regression models. *Manuscript*. (Available from <http://hedibert.org/wpcontent/uploads/2018/06/uribe-lobes-Sep2017.pdf>.)
- Wang, H., Yu, H., Hoy, M., Dauwels, J. and Wang, H. (2016) Variational Bayesian dynamic compressive sensing. In *Proc. Int. Symp. Information Theory*, pp. 1421–1425. New York: Institute of Electrical and Electronics Engineers.
- Zhu, B. and Dunson, D. B. (2013) Locally adaptive Bayes nonparametric regression via nested Gaussian processes. *J. Am. Statist. Ass.*, **108**, 1445–1456.
- Ziniel, J. and Schniter, P. (2013) Dynamic compressive sensing of time-varying signals via approximate message passing. *IEEE Trans. Signal Process.*, **61**, 5270–5284.

Supporting information

Additional ‘supporting information’ may be found in the on-line version of this article:

‘Supplement to “Dynamic shrinkage processes”’.