

Flexible Bayesian MIDAS: time-variation, group-shrinkage and sparsity*

David Kohns[†]

Galina Potjagailo[‡]

Aalto University

Bank of England

November 15, 2023

See here for most recent version!

Abstract

We propose a mixed-frequency regression prediction approach that models a time-varying trend, stochastic volatility and fat-tails in the variable of interest. The coefficients of high-frequency indicators are regularized via a shrinkage prior that accounts for the grouping structure and within-group correlation among lags. A new sparsification algorithm on the posterior motivated by Bayesian decision theory derives inclusion probabilities over lag groups, thus making the results easy to communicate without imposing sparsity a priori. An empirical application on nowcasting UK GDP growth suggests that group-shrinkage in combination with the time-varying components substantially increases nowcasting performance by reading signals from an economically meaningful sub-set of indicators, whereas the time-varying components help by allowing the model to switch between indicators. Over the data release cycle, signals initially stem from survey data and then shift towards few “hard” real activity indicators. During the Covid-19 pandemic, the model performs relatively well since it shifts towards indicators for the service and housing sectors that capture the disruptions from economic lockdowns.

Keywords: Bayesian MIDAS regressions, Forecasting, Time-variation and fat tails, Grouped Horseshoe Prior, Decision Analysis.

JEL Codes: C11, C32, C44, C53, E37

*We are grateful for highly useful comments and suggestions from Arnab Bhattacharjee, Todd Clark, George Kapetanios, Gary Koop, Silvia Miranda-Agrippino, Andre Moreira, Tatevik Sekhposyan, and the participants of the 42nd International Symposium on Forecasting, the Bergamo BWES 2022 Workshop, the Conference on Real-Time Data Analysis, Methods, and Applications at the Cleveland Fed, the Orebro 5th Annual Workshop on Financial Econometrics, the 2023 SNDE Annual Symposium, and the RES 2023 Annual Conference. *Disclaimer:* The views expressed in this paper are solely those of the authors and do not necessarily represent those of the Bank of England or the European Central Bank.

[†]Corresponding author. Email: david.kohns94@googlemail.com

[‡]Email: galina.potjagailo@bankofengland.co.uk

1 Introduction

The large and unprecedented economic fluctuations caused by the Covid-19 pandemic have reemphasized the challenge for time series and prediction models to be flexible enough to extract heterogeneous signals and to rely on different sets of indicators over time, as well as to accommodate time variation and extreme observations (Lenza and Primiceri, 2022; Carriero et al., 2021; Antolin-Diaz et al., 2021). Timely survey indicators that were previously heavily relied upon for nowcasting were less able to capture the magnitude of the economic decline and the heterogeneous impact across economic sectors (Bank of England, 2020). Even as the world has emerged from the pandemic, this remains relevant since economists increasingly monitor heterogeneous types of new data sets, and need to incorporate large economic fluctuations into their models such as the recent surge in inflation rates in many economies.

We propose a Bayesian mixed-frequency regression model that aims to address these challenges. The T-SV-t BMIDAS models gradual trend shifts and potentially extreme variance shifts in the variable of interest via a time-varying trend, stochastic volatility and t-distributed errors that can assume fat tails. If left unmodeled, such variation in the variable of interest can blur nowcast or forecast signals. Information from higher-frequency indicators enters via a multivariate mixed-data sampling (MIDAS) regression that links the variable of interest to higher-frequency lags of each indicator via parametric functional constraints (Ghysels et al., 2007, 2020). The high-frequency lag coefficients are regularized via a flexible three-tiered group-wise shrinkage prior that jointly shrinks groups of lags of less informative indicators while accounting for the correlation among lags. Finally, we propose a sparsification algorithm for the posterior motivated by Bayesian decision theory which achieves ex-post variable selection with the aim of intuitive signal communication.

For group-shrinkage, we employ the GIGG (Group Inverse-Gamma Gamma) global-local prior, where two heavy-tailed processes shrink irrelevant signals towards zero while avoiding excessive shrinkage of relevant signals globally and locally (Polson and Scott, 2010). Proposed by Boss et al. (2021) for panel settings, the prior adds a group-wise shrinkage dimension that simultaneously shrinks *between* and *within groups* of covariates. We adopt the prior for the time-series context and set hyperparameters such that between-group shrinkage with high correlation within group is preferred. This regularizes variance inflation commonly observed in over-parameterized multivariate MIDAS models (Carriero et al., 2015) and accounts for the temporal lag correlation between higher-frequency observations of the same indicator (which we define as “groups”).¹

¹ Intra-group collinearity is present in unrestricted MIDAS (Forni et al., 2015), and is reinforced when polynomial restrictions are imposed to achieve parsimony and a meaningful lag profile (Ferrara et al., 2022). Generally, three-tiered group-shrinkage can be relevant in a broader range of econometric problems where a group-structure exist, and the prior can accommodate different types of group structures.

The prior shrinks in a continuous way without imposing exact zero coefficients, which makes it more flexible. However, a drawback is that communication of the signals that the model exploits is difficult without a clear-cut selection of indicators. To address this, we propose a new sparsification algorithm motivated by Bayesian decision theory that is applied to the posterior of the MIDAS coefficients. The algorithm selects *ex-post* those high-frequency lag groups that best summarize the predictions of the model, in the spirit of Hahn and Carvalho (2015), and sets the coefficients of others to exact zeros. Imposing this via the prior directly is known to be sensitive to the hyperparameters (Barbieri and Berger, 2004; O’Hara and Sillanpää, 2009) and to the correlation structure in the data (Barbieri et al., 2021). By asking which subsets of data reproduce the posterior predictions as closely as possible, the *ex-post* sparsification step allows to assign inclusion probabilities to coefficient groups and hence to enhance the interpretability and tractability of predictions.

In an empirical application, we nowcast quarter-on-quarter GDP growth in the United Kingdom using monthly macroeconomic indicators. We conduct nowcasts in a pseudo-real-time setting over the sample period 1999 to 2021, also examining a pre-pandemic sample. We find that the proposed model performs competitively against a range of alternatives for both sample periods. During the pandemic quarters, the model detects the initial economic trough earlier and nowcasts the economic recovery more precisely.

To understand the role of the main model features, we compare mean and density nowcast performance to models that shut some or all of the time-varying components down, as well as against a range of alternative priors on the MIDAS coefficients. We find that allowing for a time-varying trend and SV with t-distributed errors, substantially improves nowcast performance when combined with the group-shrinkage prior. Examining inclusion probability patterns, we show that the GIGG prior identifies a small set of indicators with high inclusion probability, while still reading signals from other indicators albeit with lower probability and higher uncertainty. Over the data release cycle, the model reads signals from survey indicators early on and then shifts towards signals from a few real activity indicators. Moreover, we find that the combination of the time-varying components along with group-shrinkage helps exploiting the high-frequency data effectively as they get released. During the Covid-19 pandemic, the inclusion of t-distributed errors helps to channel the group-shrinkage towards relying more heavily on indicators that capture sentiment and activity in the service sector, as well as indicators on housing, both of which reflected disruptions from economic lockdowns. By contrast, models with other priors profit less from including the time-varying components. The horseshoe prior reads information in a dense and diffuse way across indicators and nowcast periods, and the spike-and-slab prior select sub-sets of indicators but shifts less effectively between them.

The proposed model combines various features and nests existing models in the literature. MIDAS regressions are popular since they are computationally less demanding than mixed-

frequency state space representations (Bai et al., 2013), and are able to exploit heterogeneous signals from indicators since they do not rely on their co-movement. Non-linear and non-parametric MIDAS structures have been used to address non-linearities in the link between indicators and target (Guérin and Marcellino, 2013; Ghysels et al., 2020). Instead, we focus on modeling non-linear features in the target variable via a time-varying trend and stochastic volatilities that account for extreme observations, since these have been found to improve predictive performance in a range of unobserved component models, VARs, and dynamic factor models (Stock and Watson, 2009; Clark, 2011; D’Agostino et al., 2013; Berger et al., 2016; Carriero et al., 2016; Antolin-Diaz et al., 2017). Our model is closest to Carriero et al. (2015) who combine a MIDAS structure with stochastic volatility, but standard Minnesota type shrinkage. Antolin-Diaz et al. (2021) allow for a time-varying trend and stochastic volatility, exploit higher-frequency indicators via a dynamic factor model, and model outliers via an additive component in GDP growth. Instead, we adopt a Student’s t-distribution to the shock structure, which can incorporate both smaller and more frequent variance shifts as well as rare extreme outliers (Jacquier et al., 2004; Clark and Ravazzolo, 2015).

The GIGG prior can be seen as an extension of the popular horseshoe prior (Carvalho et al., 2010) towards three-tiered group shrinkage. Other grouped shrinkage priors, such as Laplace-type (Casella et al., 2010) or Cauchy regularization (Xu et al., 2016), assume non-exchangeability in a group, but apply a uniform level of shrinkage within group, with no inference on the correlation via a covariate level scale. Mogliani and Simoni (2021) show that by augmenting Laplace-type regularization to be adaptive with spike-and-slab variable may provide improvements over the former priors in grouped MIDAS regressions. Xu and Ghosh (2015), on the other hand, propose a group-sparse spike-and-slab prior which applies selection both on the group and within-group level, akin to the sparse group-lasso of Simon and Tibshirani (2012). Babii et al. (2022) study this estimator from a frequentist point of view in the context of MIDAS models and find that it provides substantial improvements for predictions. The GIGG prior can, due the flexibility of its three tiers, be thought of as a continuous generalization to these approaches that can flexibly adapt to any correlation structure and is therefore suitable for any basis of the MIDAS weighting function.

Finally, the ex-post sparsification approach relates to the debate of shrinkage versus sparsity initiated in Giannone et al. (2021) who highlight the importance of separately modeling shrinkage and sparsity. Our proposed methodology, too, separates shrinkage from the prior and sparsity from group selection and is able to provide uncertainty quantification in selection via inclusion probabilities over time. A key difference to Giannone et al. (2021) or group selection methods presented in Mogliani and Simoni (2021) and Babii et al. (2022), is that we view the proposed algorithm as a tool to understand and communicate the predictions of the model, coherent with decision theory, rather than to detect sparsity as a function of the actual obser-

variations of the target variable. At the same time, the GIGG prior can be tuned to incorporate the knowledge that macroeconomic data is highly correlated which facilitates inference when many small coefficients are likely, avoiding an “illusion of sparsity” (Giannone et al., 2021) that might otherwise be present with selection priors.

The remainder of the paper is structured as follows. Section 2 presents the T-SV-t BMIDAS, the GIGG prior and sparsification step. Section 3 outlines the data set and setup of the empirical application. Section 4 discusses nowcast evaluation results and the role of the model features, and section 5 discusses the results through the lens of variable inclusion probabilities. Section 6 concludes. The accompanying appendix shows details on methodology and additional results.

2 T-SV-t BMIDAS, group-shrinkage and sparsification

2.1 BMIDAS with time-varying components

The proposed T-SV-t BMIDAS combines time-varying component features with a multivariate MIDAS regression: (1) a time-varying trend component (2) stochastic volatility processes, and (3) t-distributed errors in the the observation equation that allow for fat tails. Let y_t be the scalar valued target variable observed at discrete time points $t = (1, \dots, T)$ which is modelled by unobserved states at frequency t and K covariates intermittently observed $m \geq 1$ times, $x_{t-(j-1)/m,k}$ for $j = (1, \dots, m)$ and $k = (1, \dots, K)$. In this paper, t will refer to quarters and $m = 3$ to months. Note that for ease of exposition, we choose m to be the same for all k , which will be relaxed following section 3. The model takes the following state-space form:

$$y_t = \tau_t + \sum_{k=1}^K \mathcal{B}(L^{1/m}; \theta_k) x_{k,t} + \sqrt{\lambda_t} e^{\frac{1}{2}(h_0 + w_h \tilde{h}_t)} \tilde{\epsilon}_t^y, \quad (1)$$

$$\tilde{\epsilon}_t^y \sim N(0, 1), \quad \lambda_t \sim IG(\nu/2, \nu/2)$$

$$\tau_t = \tau_{t-1} + e^{\frac{1}{2}(g_0 + w_g \tilde{g}_t)} \tilde{\epsilon}_t^\tau, \quad \tilde{\epsilon}_t^\tau \sim N(0, 1) \quad (2)$$

$$\tilde{h}_t = \tilde{h}_{t-1} + \tilde{\epsilon}_t^h, \quad \tilde{\epsilon}_t^h \sim N(0, 1)$$

$$\tilde{g}_t = \tilde{g}_{t-1} + \tilde{\epsilon}_t^g, \quad \tilde{\epsilon}_t^g \sim N(0, 1), \quad (3)$$

where (1) is the observation equation and (2)-(3) describe the evolution of states, latent trend and stochastic volatilities, respectively.

The trend component τ_t is modelled as a driftless random walk. This has long tradition in inflation trend estimation (see e.g. Stock and Watson (2007); Harvey et al. (2007)), but has recently also been introduced to nowcast models of GDP growth using dynamic factor models (Antolin-Diaz et al., 2017, 2021). The rationale for a unit root process in the trend,

compared to the assumption of discrete changes in mean via structural breaks as in Kim and Nelson (1999) or McConnell and Perez-Quiros (2000) is to capture slow moving changes in GDP growth that may not necessarily be captured by the higher frequency components. This information typically relates to short run fluctuations of the business-cycle.² Nevertheless, which type of fluctuation the trend captures is not unequivocal, as the degree of smoothness of the latent trend depends on the priors employed and on the residual variation explained by the other model components (1) and (3). In the empirical application, we therefore assess the role of the trend by shutting down model components, and check the sensitivity of the trend to model specifications and priors.

Conditionally on the trend, a MIDAS component, $\sum_{k=1}^K \mathcal{B}(L^{1/m}; \theta_k)x_{k,t}$, relates short-term fluctuations from a possibly wide range of high-frequency indicators to y_t . Define $L^{1/m}$ as the back-shift operator such that $L^{1/m}x_{t,k} = x_{t-1/m,k}$, and denote by θ a L -dimensional vector of coefficients. Then $\mathcal{B}(L^{1/m}, \theta_k)$ defines a lag polynomial, which, as is common in the MIDAS literature (e.g. Babii et al. (2022)), is parameterized by a weighting function, $\omega : \mathbb{R} \times \mathbb{R}^L \rightarrow \mathbb{R}$:

$$B(L^{1/m}, \theta_k)x_{t,k} = \sum_{j=1}^{qm} \omega\left(\frac{j-1}{qm}; \theta_k\right)x_{t-(j-1)/qm,k}. \quad (4)$$

Here, q indicates how many quarterly lags enter the polynomial in addition to the contemporaneous m months. Many different weighting functions have been proposed in the literature (Ghysels et al., 2007). In this paper, we consider ω to be a linear functional, $\omega(\theta_k) = \sum_{l=1}^L \theta_{k,l} \omega_{l,j}$ for $l = (1, \dots, L)$ of which unrestricted weights (U-MIDAS, see Ghysels et al. (2007); Foroni et al. (2015)), and restricted Almon lag polynomial weights (Almon, 1965; Ferrara et al., 2022) are popular subsets thereof³⁴. While U-MIDAS treats each higher frequency lag polynomial as an independent covariate, Almon polynomials force the curvature of regression coefficients belonging to a given covariate k to adhere to a L -degree polynomial process. This has particular appeal for macro time-series applications since economically relevant restrictions to the polynomial's end-points can be enforced, such as constraining the functional to peter out smoothly to zero for distant lags (Smith and Giles, 1976). Denote these restrictions by $r \forall k$. While the restrictions may induce parsimony when $L - r \ll qm$, the implied linear transformations typically induce a highly collinear grouping structure that can cause mixing problems for shrinkage priors (Griffin and Brown, 2012). The importance of modelling this correlation will become

² A different non-stochastic approach to account for long-run growth features is presented in Giannone et al. (2019), where iterative forecasts are enforced to return to a long-run cointegrating equilibrium.

³ Foroni and Marcellino (2014) show that linear MIDAS methods are competitive with non-linear MIDAS weighting schemes such as the non-linear Almon and beta functions (Ghysels et al., 2004, 2007; Andreou et al., 2010; Ghysels et al., 2020), while being compatible with off the shelf shrinkage methods. Linear MIDAS have been employed in (Foroni et al., 2015; Angelini et al., 2011; Mogliani and Simoni, 2021).

⁴ In general, however, any arbitrary function can be used for ω which Babii et al. (2022) denote as the dictionary of functions.

apparent when comparing priors that assume within lag group exchangeability to the GIGG prior and its sparsification.

Finally, we allow for two types of residual variation: (1) stochastic volatilities, $\{\tilde{h}_t, \tilde{g}_t\}$, for GDP growth, y_t , and the latent trend, τ_t , that model persistence in error variance, and (2) student-t errors in the y_t equation, to model non-persistent shocks such as extreme events and outliers that are otherwise unaccounted for by the other model components (Jacquier et al., 2004; Clark and Ravazzolo, 2015). Since states $\{h_t, g_t\}$ can be regarded as parameters, the state components alone can quickly make the model high dimensional. To allow for stronger shrinkage on the states via priors on w_h and w_g , which control state-smoothness, we use the non-centered state space formulation after Frühwirth-Schnatter and Wagner (2010). This retains conditional conjugacy with the priors discussed below.⁵ The introduction of t-errors may benefit density nowcasts since fat-tails imply wider probability bands around extreme observations, but can also be relevant for point nowcasts because t-distributed errors discount large contemporaneous movements in y_t and, therefore, limit the propagation of outliers to the posteriors of the model components (Chiu et al., 2017). Differently from linear outlier treatments on the level of the target variable (Stock and Watson, 2016; Antolin-Diaz et al., 2021), t-distributed errors reflect a continuous range of time variation and can incorporate both smaller, more frequent variance shifts and rare extreme outliers.

2.2 Priors

We now describe the prior hierarchy used for drawing inference on all unknowns in (1)-(3):

$$\pi(\zeta) = \pi(\theta)\pi(\boldsymbol{\tau})\pi(\tilde{\mathbf{h}})\pi(\tilde{\mathbf{g}})\pi(\phi)\pi(\boldsymbol{\lambda}|\nu)\pi(\nu), \quad (5)$$

where ζ collects all unknowns into one vector. The MIDAS coefficients θ are regularised by the GIGG prior with group-shrinkage structure. Bold-faced letters refer to time-ordered vectors (e.g., $\boldsymbol{\tau} = (\tau_1, \dots, \tau_T)'$), and ϕ collects any remaining state parameters $(\tau_0, h_0, g_0, w_h, w_g)'$. $\{\tau_0, h_0, g_0\}$ are the starting values for the latent states.

2.2.1 GIGG prior on MIDAS coefficients

Multivariate MIDAS regressions are highly parameterised, but the grouping and correlation structure across coefficients can be exploited for efficient shrinkage. We do so in three ways:

⁵ Since w_g and w_h appear in the observation and state equation, one can apply conjugate normal priors which exert stronger shrinkage than the inverse-gamma priors conventionally applied in Bayesian state-space models. This approach has been used in large time-varying parameter models, see Huber et al. (2021). The more commonly employed centered SV process h_t can be exactly recovered given that $h_t = h_0 + w_h \tilde{h}_t$ (likewise for g_t).

shrinkage across all coefficients, group-wise shrinkage jointly for coefficients belonging to an indicator, and modelling of the correlation between coefficients within a group. The prior, therefore, speaks to the fact that, first, not all of the k indicators are equally relevant for prediction, addressed the *group-wise shrinkage* of lags belonging to a given indicator. Second, correlation among parameters *within the lag group* is typically very high for consecutive high-frequency lags, particularly when functional restrictions such as those implied by the Almon polynomial are applied. Exchangeable priors that leave this unaddressed, are liable to random covariate selection and bad mixing (Boss et al., 2021; Piironen et al., 2020). Third, group-structure and correlation across lags can *jointly* matter for predictive performance since the relative impact of the lag group on the target variable can be affected. The prior, “inverse-Gamma Gamma” (GIGG) (Boss et al., 2021), is specified as

$$\begin{aligned} \theta_{k,j} &\sim N(0, \vartheta^2 \gamma_k^2 \varphi_{k,j}^2), \quad \forall j \in \{1, \dots, L-r\} \\ \vartheta &\sim C_+(0, 1), \quad \gamma_k^2 | a_k \sim G(a_k, 1), \quad \varphi_{k,j}^2 \sim IG(b_k, 1), \end{aligned} \tag{6}$$

where $G(\bullet, \bullet)$, $IG(\bullet, \bullet)$ and $C_+(\bullet, \bullet)$ refer to the Gamma, inverse-Gamma, and half Cauchy distribution with positive support, respectively. ϑ controls the overall level of sparsity. γ_k acts as a shrinkage factor on the joint impact of coefficients from group k , and $\varphi_{k,j}$ controls how correlated group members are within k . The magnitude and relative size of a_k and b_k determine the amount of shrinkage and relative importance between group-shrinkage and within-group correlation.

Depending on the choice of hyperparameters, prior (6) is related to several previous approaches employed in MIDAS settings. For a group-size of 1 and setting $a_k = b_k = 0.5$, it reduces to the horseshoe prior of Carvalho et al. (2010) which Kohns and Bhattacharjee (2022) apply to nowcasting US GDP growth with U-MIDAS sampled components. The GIGG can also be viewed as a continuous generalisation to the group adaptive lasso with spike-and-slab group selection which Mogliani and Simoni (2021) analyse empirically and theoretically for Almon MIDAS regressions. Instead of Laplacian slab distributions, however, the GIGG allows for fatter tailed distributions and inference on the correlation of coefficients within a group. More broadly, for any group-size larger than 1, the prior follows a correlated normal beta prime distribution akin to a grouped formulation of the prior by Armagan et al. (2013), $\gamma_k^2 \varphi_{k,j}^2 \sim \beta'(a_k, b_k)$ which flexibly nests the group-horseshoe Xu et al. (2016).⁶ and may also be viewed as a continuous approximation to the group-sparse spike-and-slab prior of Xu and Ghosh (2015).

For the application to nowcasting, we leverage the model structure to inform on the hyperparameters of the GIGG prior. Generally, the lower a_k , the stronger the group is penalised, while the higher b_k , the more within-group correlation is allowed for. But the relative magnitudes

⁶ Unlike Xu et al. (2016), the prior reduces to the exact horseshoe at group-size 1.

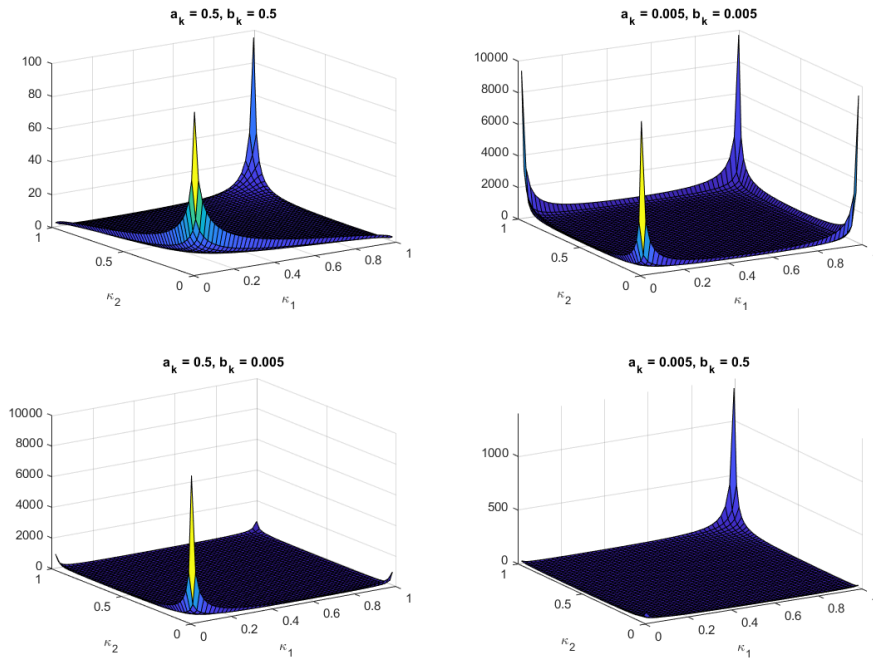


Figure 1: Bi-variate shrinkage coefficient plots for various hyper-parameter values. a_k controls group-level shrinkage while b_k controls the degree of correlation within groups.

matter as well. To visualise these relationships, Figure 1 shows the prior distribution on the shrinkage coefficients $\kappa_j \in [0, 1]$, implicitly defined by $(\vartheta, \gamma_k, \varphi_{k,j})$, for different choices for a_k and b_k and a group of 2. These are often analysed in the context of generalised linear models (Piironen et al., 2017) since they inform on how much the prior shrinks toward a mean of 0 from the maximum likelihood estimate of the regression weights (see appendix A.1.1 for full derivations the MIDAS component). High probability mass on 1 implies aggressive shrinkage toward 0 and the degree of concentration of the joint probability distribution informs on the degree of correlation in shrinkage.

Two general shrinkage profiles emerge from choosing $a_k = b_k$ or $a_k \neq b_k$. Setting $a_k = b_k$ induces a horseshoe type shape, which displays high correlation when set to $a_k = b_k = 0.5$ (upper left plot). The high degree of correlation favours to either jointly shrink the group members to 0, or leave the coefficients relatively un-shrunk, akin to the group-horseshoe by Xu et al. (2016). For smaller values such as 0.005 (upper right plot), on the other hand, we observe the tendency for independent horseshoe type shrinkage, implying low correlation. A low magnitude of a_k and b_k can be appropriate when prior knowledge exists that only selected lags are important. When $a \neq b$, within-group correlation is high and the horseshoe shape collapses towards very high or very little joint shrinkage for *both* coefficients. Choosing $a \gg b$ (lower left sub-plot) favours little regularisation (high mass at South edge). Choosing $a \ll b$ (lower right

sub-plot) induces aggressive joint shrinkage (high mass at North edge), and the overall impact of both coefficients is pushed toward sparsity. Nonetheless, in both cases, small mass in the edges retains the flexibility to escape the joint shrinkage tendencies.

For our empirical application, we adopt this latter case of aggressive group-correlated shrinkage, setting $a = 1/T, b = 0.5$ for all coefficient groups. This choice reflects the view that a high degree of correlation is present between MIDAS lag coefficients belonging to the same indicator, particularly when coefficients are restricted via Almon polynomials. And a comparison across different hyperparameter choices suggests an improved nowcasting performance with strong prior correlation in shrinkage, particularly prior to the pandemic.

2.2.2 Priors for the latent states and other coefficients

For priors related to the states and starting conditions, we follow the previous literature. For the latent states $(\boldsymbol{\tau}, \tilde{\mathbf{h}}, \tilde{\mathbf{g}})$ we consider a joint normal prior derived using methods proposed in Chan and Jeliazkov (2009) and McCausland et al. (2011). This allows representing the entire conditional state posterior as a tractable normal distribution. Latent states are prone to overfitting in over-parameterised models and the smoothness of the state processes are notoriously difficult to identify from the data alone (Cogley et al., 2005; Antolin-Diaz et al., 2017). We therefore put normal priors on the state standard deviations with small prior variance that exert stronger shrinkage than commonly employed inverse-Gamma priors (Frühwirth-Schnatter and Wagner, 2010), $w_i \sim N(0, 0.01^2)$ for $i \in \{h, g\}$. This encourages the states to be slow moving. We set weakly informative priors on the initial conditions $\{i_0\}_{i \in \{h, g, \tau\}} \sim N(0, 10)$. Lastly, we put a relatively uninformative uniform prior on the degrees of freedom of the t-distributed errors in the observation equation, $\nu \sim \mathcal{U}[2, 50]$. The lower bound is chosen to insure the existence of the first two moments of the t-distribution (Chan and Grant, 2016).

2.3 Estimation Algorithm

Here, we briefly review the estimation algorithm, firstly defining the likelihood stacked likelihood across all observation. Let $\mathbf{1}_T$ be a column of ones. Recall that $\mathbf{h} = h_0 \mathbf{1}_T + w_h \tilde{\mathbf{h}}$ and $\mathbf{g} = g_0 \mathbf{1}_T + w_g \tilde{\mathbf{g}}$ and define $\boldsymbol{\tau} = (\tau_1, \dots, \tau_T)$. Then it follows that

$$(\mathbf{y}|\mathbf{X}, \theta, \boldsymbol{\tau}, \tilde{\mathbf{h}}, h_0, w_h, \boldsymbol{\lambda}) \sim N(\mathbf{Z}\boldsymbol{\theta} + \boldsymbol{\tau}, \Lambda_h \Lambda_\lambda), \quad (7)$$

where $\mathbf{Z}\boldsymbol{\theta}$ are the vectorised MIDAS coefficients and $\Lambda_h = \text{diag}(e^{h_1}, \dots, e^{h_T})$ and $\Lambda_\lambda = \text{diag}(\lambda_1, \dots, \lambda_T)$.⁷ The assumption of independence between MIDAS coefficients and state

⁷ Define $X_k = (x_{k,t-\frac{j-1}{qm}})_{t \in \{1, \dots, T\}, j \in \{1, \dots, qm\}}$ as a $T \times qm$ matrix of high frequency lags $\forall k$ and $W = (w'_1, \dots, w'_{L-r})$ as a $qm \times L-r$ matrix of Almon restricted weights. Then, $Z_k = X_k W$ and the stacked

space components allows to derive conditional posterior distributions via an efficient Metropolis-within-Gibbs sampling algorithm. We sequentially simulate from the following algorithm:

1. Sample $\theta \sim p(\theta|\mathbf{y}, \bullet)$
2. Sample hyper-parameters $(\vartheta, \gamma_k^2, \varphi_{k,j}^2, \nu_\vartheta)$ in one block
 - (a) $\vartheta^2 \sim p(\vartheta^2|\mathbf{y}, \bullet)$
 - (b) $\gamma_k^2 \sim 1/p(\gamma_k^{-2}|\mathbf{y}, \bullet)$ for $k = (1, \dots, K)$
 - (c) $\varphi_{k,j}^2 \sim p(\varphi_{k,j}^2|\mathbf{y}, \bullet)$ for $j = (1, \dots, L - r)$
 - (d) $\nu_\vartheta \sim p(\nu_\vartheta|\mathbf{y}, \bullet)$
3. Sample $\tilde{\tau} \sim p(\tilde{\tau}|\mathbf{y}, \bullet)$ and $\tau_0 \sim p(\tau_0|\mathbf{y}, \bullet)$
4. Sample $\tilde{\mathbf{h}} \sim p(\tilde{\mathbf{h}}|\mathbf{y}, \bullet)$, $h_0 \sim p(h_0|\mathbf{y}, \bullet)$ and $w_h \sim p(w_h|\mathbf{y}, \bullet)$
5. Sample $\tilde{\mathbf{g}} \sim p(\tilde{\mathbf{g}}|\mathbf{y}, \bullet)$, $g_0 \sim p(h_0|\mathbf{y}, \bullet)$ and $w_g \sim p(w_g|\mathbf{y}, \bullet)$
6. Sample $\lambda_t \sim p(\lambda_t|\mathbf{y}, \bullet)$ for $t = (1, \dots, T)$
7. Sample $\nu \sim p(\nu|\mathbf{y}, \bullet)$ with a Metropolis step.

ν_ϑ is an $IG(\bullet, \bullet)$ distributed mixture variable that, when $\vartheta^2|\nu_\vartheta \sim IG(1/2, 1/\nu_\vartheta)$ and $\nu_\vartheta \sim IG(1/2, 1)$, gives exact draws for $(\vartheta|\mathbf{y}, \bullet)$ (Makalic and Schmidt, 2015). We iterate sampling steps 1.-7. 5000 times for burn-in and retain further 5000 samples for inference⁸. To speed up the computations of the posteriors, we make use of the state sampling techniques of Chan and Jeliazkov (2009) and McCausland et al. (2011) for the state variables. This allows drawing steps 3.-5. in a non-recursive fashion, which increases sampling and computational efficiency via sparse-matrix operations compared to Kalman filter based techniques (Carter and Kohn, 1994). We sample from the posteriors of $(\tilde{\mathbf{h}}, \tilde{\mathbf{g}})$ using the approximate sampler of Kim et al. (1998). The conditional MIDAS coefficient posterior is normal and we use the algorithm of Bhattacharya et al. (2016) to speed up computation and aid mixing when $\sum_{k=1}^K (L - r) \gg T$. See appendix A.1 for details on the conditional posterior distributions and the Metropolis step.

2.4 Group-sparsification on the posterior

Continuous priors such as the GIGG may already achieve optimal empirical and prediction risk properties (Chakraborty et al., 2020; Boss et al., 2021). However, when the policy maker aims to understand and communicate nowcasts, the fact that with such priors, posteriors remain non-zero can be a limitation since the impact on nowcasts remains opaque. In view of this, we propose to *ex-post* sparsify the posterior of the regression weights based on decision theory, with the aim of finding the smallest subset of indicators which achieve closest predictive performance

matrix is $\mathbf{Z} = (Z_1, \dots, Z_T)$ and the MIDAS functional coefficient vector is $\theta = (\theta'_1, \dots, \theta'_K)'$.

⁸ To check convergence, we tested with 20000 iterations for burn-in and 20000 for inference, and results remain similar.

to the unsparsified model. Compared to priors that conduct selection, we view sparsity as a decision tool that is separate from any regularisation that the prior imposes. Giannone et al. (2021) highlight the importance of separating shrinkage from sparsity, particularly for macroeconomic data, where high correlation among indicators can result in falsely detecting sparsity (“illusion of sparsity”).

We derive a new analytical solution for ex-post sparsification for the general MIDAS posterior which takes the temporal grouping structure of the MIDAS functional into account. Specifically, we threshold the coefficients of lags pertaining of an indicator jointly to zero if they have little effect on the predictions. And we use insights from the sparse group lasso literature (Simon and Tibshirani, 2012; Breheny and Huang, 2015) to formulate a computationally efficient solution to the group selection problem which is robust to the intra-group correlation, and applicable to any MIDAS specification using linear weights for (4).

Suppose we take the perspective of a decision maker who minimises a utility function over the Euclidean distance between a linear model that penalises group-size akin to Zou (2006) and the predictions from model (1)-(3).

$$\mathcal{L}(\tilde{\mathbf{Y}}, \alpha) = \frac{1}{2} \|\mathbf{Z}\alpha - \tilde{\mathbf{Y}}\|_2^2 + \sum_{k=1}^K \phi_k \|\alpha_k\|_2, \quad (8)$$

where $\tilde{\mathbf{Y}}$ refers to a realisation from the posterior predictive distribution $p(\tilde{\mathbf{Y}}|\mathbf{y}) = \int p(\tilde{\mathbf{Y}}|\mathbf{y}, \theta, \bullet) p(\theta|\mathbf{y}, \bullet) d\theta$, and $\|\bullet\|_p$ refers to the ℓ_p -norm.⁹ α is an unknown MIDAS weight vector which due to the group-wise penalty $\sum_{k=1}^K \phi_k \|\alpha_k\|_2$ may display sparsity via penalty parameter ϕ_k for each indicator group. Similar to the logic of adaptive group-lasso (Wang and Leng, 2008), the penalisation term induces non-differentiability at zero, which creates a soft-thresholding effect between $[-\phi_k, \phi_k]$, thereby forcing the coefficients on all group members to zero. Unlike traditional selection approaches such as the spike-and-slab prior, sparsity is induced using the prediction of $\mathbf{Y}$ as a target, rather than $\mathbf{Y}$ directly which has been previously shown to benefit the stability of subset selection (Piironen and Vehtari, 2017; Piironen et al., 2020).

The Bayes optimal solution for α is obtained by integrating out the two sources of uncertainty which characterise the uncertainty in $\tilde{\mathbf{Y}}$: posterior uncertainty from the predictive distribution, as well as in the parameters θ (Lindley, 1968). Here, we show the analytical solution for (8) and discuss the assumptions needed to derive it (for a full derivation, see appendix A.2). The sparsified estimate $\alpha_k^{*(s)}$ for each Gibbs-sampling step $s = 1, \dots, S$, is given by:

$$\alpha_k^{*(s)} = \left(\|\theta_k^{(s)}\|_2 - \phi_k^{(s)} \right)_+ \frac{\theta_k^{(s)}}{\|\theta_k^{(s)}\|_2}, \quad \forall k \in \{1, \dots, K\}, \quad (9)$$

⁹ Note that for simplicity we define the predictive distribution over in-sample values of $\mathbf{Z}$, but in principle any data can be used for the analysis.

where $(x)_+ = \max(x, 0)$. (9) implies that when $\theta_k^{(s)}$ are close to $\mathbf{0}$, then $\alpha_k^{*(s)} = \mathbf{0}$, whereas, when $\theta_k^{(s)}$ are sufficiently large, then $\alpha_k^{*(s)} = (1 - \frac{\phi_k^{(s)}}{\|\theta_k^{(s)}\|_2})\theta_k^{(s)}$, in which case the first term will be very close to 1, thus imposing close to no further shrinkage. Two assumptions are needed to derive (9). Firstly, it requires orthonormalisation of the data for each k such that $T^{-1}\tilde{\mathbf{Z}}_k'\tilde{\mathbf{Z}}_k = I_{L-r}$. This serves to simplify the make the solution robust to any correlation implied by ω within (4). Further, as shown in Simon and Tibshirani (2012), not orthonormalising groups ignores the cross-correlation of group members in k , in which case the algorithm implicitly prefers to not threshold groups with large covariance, and ignores that $\mathbf{Z}_k$ might have different scales.¹⁰ Secondly, we make use of work by Ray and Bhattacharya (2018) and Chakraborty et al. (2020) who show that, when setting $\phi_k^{(s)} = \frac{1}{\|\theta_k^{(s)}\|_2}$, iterative solution methods such as the coordinate descent (Friedman et al., 2010) converge after the first cycle.

We then quantify (9) the posterior “inclusion probability” that a given group of coefficients was included in the model’s forecasts:

$$p(\alpha_k^{*(s)} \neq 0 | \tilde{\mathbf{Y}}) \approx \frac{1}{S} \sum_{s=1}^S \mathbb{1}_{\alpha_k^{*(s)} \neq 0} \quad \forall k, \quad (10)$$

and associated posterior variance: $\text{var}(\alpha_k^{*(s)} \neq 0 | \tilde{\mathbf{Y}}) \approx \frac{1}{S} \sum_{s=1}^S (\mathbb{1}_{\alpha_k^{*(s)} \neq 0} - \bar{\mathbb{1}}_{\alpha_k \neq 0})^2$. Thus, posterior inclusion probabilities represent the relative frequency with which the group k was selected in the sparsified estimate $\alpha^{*(s)}$ over all Gibbs draws. These reflect the uncertainty around the inclusion of a variable and its relative impact, which would not be the case if we were to sparsify only once on the posterior mean.¹¹

Linking regularisation from priors and variable selection has long tradition in statistics (Lindley, 1968). Hahn and Carvalho (2015) re-popularised this approach for high-dimensional problems by introducing independent lasso-type penalties into the decision maker’s utility function. Ray and Bhattacharya (2018) and Chakraborty et al. (2020) propose computationally efficient solutions for this decision task for individual covariate and group-selection, respectively. These have been adapted to time-varying parameter models and VARs (Huber et al., 2021; Hauzenberger et al., 2021), where Huber et al. (2021) also sparsify over MCMC draws and derive posterior inclusion probabilities. But to our knowledge, existing applications to MIDAS settings have assumed independent covariate selection (Kohns and Bhattacharjee, 2022). And differently from the previous approaches, the group-sparse solution is robust to correlation

¹⁰ It can be further shown that orthonormalising the objective, establishes connection to best subset selection and uniformly most powerful invariant testing (Simon and Tibshirani, 2012)

¹¹ See Woody et al. (2021) and Chakraborty et al. (2020) for formal justification of the model selection uncertainty and the asymptotic risk properties, respectively. Note, that unlike spike-and-slab priors, we do not conduct direct inference on the inclusion probability of a variable, hence deriving its approximate posterior intervals are not immediately available.

within group and therefore generalises to any linear MIDAS approach.

We note at this point, that our sparsification approach is no panacea for detecting true sparsity of the data generating process. Which indicators are ultimately thresholded to zero, depends on the prior applied for the model and its ability to apply aggressive shrinkage to variables that are weakly related to the target, y_t . To highlight the importance for group shrinkage in MIDAS regression, we will apply our sparsification algorithm also to previously used priors for macro settings which assume exchangeability within group such as the horseshoe prior and spike-and-slab. While Giannone et al. (2021) conclude that the data are typically non-informative about the which subset of predictors to include in the model, we show that this is not necessary the case when imposing three-tiered group shrinkage and then sparsifying ex-post, which in our application helps us communicating intuitive sub-sets of most relevant predictors over time.

3 Empirical setup

For our application, we nowcast real quarter-on-quarter GDP growth in the United Kingdom based on a set of monthly macroeconomic indicators following a stylised publication calendar. In the following, we outline the data set, the stylised publication calendar we follow, and the set-up of our nowcasting exercise and evaluation.

3.1 Data set

The set of monthly macroeconomic indicators has been compiled to reflect information on the UK economy that policymakers actively monitor to gauge economic activity in real time, and is comparable to data sets employed in previous studies (Antolin-Diaz et al., 2017; Anesti et al., 2018). We include a range of real activity and survey indicators, including indices of production and services, exports and imports, a range of labour market series, as well as timely business and consumer surveys (CBI survey, PMIs, GFK). In order to capture lending conditions that can affect economic conditions via financial conditions we also include mortgage lending approvals and VISA credit card consumer spending. These series also tracked consumer spending during the pandemic, reflecting shut-downs of business and housing activity.¹² We do not add asset prices or other financial indicators which have been found to contribute little to nowcast updates

¹² Alternatively, various studies have exploited new data sources at a daily or weekly frequency (Ng, 2021; Baumeister et al., 2021; Woloszko, 2020; Huber et al., 2020; Kapetanios et al., 2022). Including such data faces the challenge of very short time-series, and we therefore abstain from including them in our main analysis. We have, however, explored the possibility of linking the monthly VISA consumer spending data with experimental weekly debit and credit card data provided by the UK’s Office for National Statistics. This analysis shows additional nowcast gains in the first weeks of the reference quarter, which however dissipate once the monthly series becomes available. Result for this exercise are available upon request.

once information from monthly survey and real activity data is accounted for (Bańbura et al., 2013; Anesti et al., 2018). Also, during the Covid-19 period financial markets were detached from real activity in the UK—asset prices initially collapsed, then stabilised early on in the pandemic in response to monetary policy interventions, and subsequently exhibited a boom that was not in line with the weakness of the real economy. The series are transformed to be approximately stationary prior to estimation.¹³

We consider the sample period from 1999Q1 to 2021Q3. The start of the sample is pinned down by data availability since many of the monthly indicators are not available for earlier years.¹⁴ To mimic incoming information over the data release cycle that a nowcaster would face in reality, we produce nowcasts based on a pseudo real time data calendar, as outlined below. However, for ease of analysis we use final vintages of the data, downloaded in December 2021. Since our focus here lies in understanding the proposed model, we leave an account for real time data releases and revisions for future research.¹⁵

3.2 Nowcast exercise

Macroeconomic data are published asynchronously at different points in time and with delays ranging from various weeks (survey data) to up to various months (labour market data) after the reference month. To simulate the information set available to the nowcaster over the data release cycle, we follow a stylised pseudo real-time data release calendar (see Table 1).

As is common with MIDAS approaches, we start the nowcast exercise for each quarter anew. We start predicting with all available information on the first of the month of the reference quarter. Following the stylised release calendar in Table 1, we generate overall 20 nowcasts that are being produced at each date in the quarter when new data series are typically released, until the release of quarterly GDP six weeks after the end of the reference quarter.¹⁶ For each new data release over the data release cycle, we generate nowcasts from the predictive distribution $p(y_{t+1}|\Omega_T^v)$, where $v = 1, \dots, 20$ refers to the nowcast periods and Ω_T^v represents the information set that expands with each data release. Since the MIDAS framework belongs to the class of reduced-form mixed frequency models, each information set Ω_T^v results in a different model, depending on which data are observable over the data release cycle (Carriero

¹³ See Table B1 in the appendix for an overview of the data and their respective transformations.

¹⁴ Some of the series have missing values at the beginning of the sample period. We interpolate these based on a principal component (PCA) model that accounts for missing information via the alternating least square algorithm, and results also remained similar when employing the Expectation Maximization algorithm (Bańbura and Modugno, 2014) instead.

¹⁵ See Anesti et al. (2018) for an analysis of UK data on the forecastability of different vintages and how to incorporate that information for nowcast updates.

¹⁶ We refer to the first GDP publication, available about 40 days after the reference quarter. We abstain from accounting for a less accurate preliminary GDP estimate that was available 25 days after the reference quarter prior to Sir Charles Bean’s 2018 review of UK economic statistics (Scruton et al., 2018).

Table 1: Stylised pseudo real-time data release calendar.

Nowcast	Quarter	Days to GDP	Month	Timing within month	Release	Publication Lag
1	Reference quarter (nowcast)	135	1	1st of month	PMIs	m-1
2		125	1	End of 2nd week	IoP, IoS, Ex, Im	m-2
3		120	1	3rd week	Labour market data	m-2
4		115	1	3rd Friday of month	Mortgage & Visa	m-1
5		110	1	End of 3rd week	CBIs & GfK	m
6		105	2	1st of month	PMIs	m-1
7		97	2	Mid of 2nd week	Quarterly GDP	q-1
8		95	2	End of 2nd week	IoP, IoS, Ex, Im	m-2
9		90	2	3rd week	Labour market data	m-2
10		85	2	3rd Friday of month	Mortgage & Visa	m-1
11		80	2	End of 3rd week	CBIs & GfK	m
12		75	3	1st of month	PMIs	m-1
13		65	3	End of 2nd week	IoP, IoS, Ex, Im	m-2
14		60	3	3rd week	Labour market data	m-2
15		55	3	3rd Friday of month	Mortgage & Visa	m-1
16		50	3	End of 3rd week	CBIs & GfK	m
17	Subsequent quarter (backcast)	45	1	1st of month	PMIs	m-1
18		35	1	End of 2nd week	IoP, IoS, Ex, Im	m-2
19		30	1	3rd week	Labour market data	m-2
20		25	1	3rd Friday of month	Mortgage & Visa	m-1

Notes: “Timing” refers to typical data release times as of December 2021, abstaining from changes in the publication calendar over the sample period. “Release” refers to the data series updated at a given nowcast, see also Table B1 in the appendix for a list of data series included. “Publication lag” represents the delay relative to the reference quarter (i.e. publication at any point in the subsequent month considered to be one month lag, m-1).

et al., 2015). To draw samples from the predictive distribution, we integrate over all parameter uncertainties which is easily implemented via Monte Carlo integration (Cogley et al., 2005).

We start the nowcast exercise with an in-sample period of 1999Q1-2011Q1, and iteratively expand it until the end of the forecast sample, $T_{end} = 2021Q3$. Since the Covid-19 pandemic represents a historic shock to the macroeconomy, we separately evaluate nowcasts over a sample that ends in 2019Q4 and one that cover the full sample period including the Covid-19 shock.

Point nowcasts are computed as the mean of the posterior predictive distribution and are compared via root-mean-squared-forecast-error (RMSFE) calculated at each nowcast period as:

$$\text{RMSFE} = \sqrt{\frac{1}{T_{end}} \sum_{t=1}^{T_{end}} (y_{T+t} - \hat{y}_{T+t|\Omega_{T+t-1}^v}^v)^2}, \quad (11)$$

where $\hat{y}_{T+t|\Omega_{T+t-1}^v}^v$ is the mean of the posterior prediction for nowcast period v using information until $T + t - 1$ and T is the initial in-sample length. Forecast density fit is measured by the mean continuous rank probability score (CRPS):

$$\text{CRPS} = \frac{1}{T_{end}} \sum_{t=1}^{T_{end}} \frac{1}{2} \left| y_{T+t} - y_{T+t|\Omega_{T+t-1}^v}^v \right| - \frac{1}{2} \left| y_{T+t|\Omega_{T+t-1}^{v,A}}^{v,A} - y_{T+t|\Omega_{T+t-1}^{v,B}}^{v,B} \right|. \quad (12)$$

with $y_{T+j|\Omega_{T+j-1}^v}^{v,A,B}$ independently drawn from the posterior predictive density $p(y_{T+1}^v|\Omega_{T+j-1}^v|y_T)$. The CRPS belongs to the class of strictly proper scoring rules (Gneiting and Raftery, 2007), and can be thought of as the probabilistic generalisation of the mean-absolute-forecast-error. The objective in terms of predictive precision is to minimise both evaluation metrics.

4 Model features and empirical nowcast performance

Next, we discuss these results further, in particular how the time-varying components and the group-shrinkage prior, respectively and jointly, contribute to the nowcast performance. In section 4.1, we compare the proposed T-SV-t-BMIDAS model with alternative models where we shut down time-variation, and in section 4.2, and we consider alternative priors on the MIDAS components for the model with and without time-varying components.

4.1 Role of time-varying unobserved components

Figure 2 looks at the role of the time-varying trend and stochastic volatilities with t-distributed errors by comparing the T-SV-t BMIDAS model with trend and stochastic volatility with t-distributed errors with the following alternatives, as well as an AR(2) benchmark (yellow line).

- *T-SV*: time-varying trend and stochastic volatility, normally distr. errors.
- *T, Const Var*: time-varying trend with constant variance.
- *SV-t*: no trend, stochastic volatility with t-distributed errors.
- *SV*: no trend, stochastic volatility, normally distributed errors.
- *Const Var*: no trend, constant variance.

Prior to the pandemic (left panels), there are substantial improvements in point and density nowcast for most models against the AR(2), and in particular for the proposed T-SV-t model compared to alternative models. Adding either a time-varying trend or stochastic volatility to the model improves point and density forecasts significantly by about 20% relative to the AR(2) benchmark, on average across nowcast periods, as shown in Table 2. Over the data release cycle, improvements of these models are only significant late in the data release cycle, however. Combining both the time-varying trend and SV-t in the proposed model leads to a stronger improvement by 35% on average, and also stabilises performance over the data release cycle, with significant improvements throughout. By contrast, the model without trend and with constant volatility shows a volatile performance across nowcast periods and does not improve against the benchmark. This underlines that incorporating at least one, and preferably both of the proposed time-varying components is important to exploit high-frequency information with BMIDAS models, in line with existing evidence based on other models for the United States (Antolin-Diaz et al., 2017; Carriero et al., 2015).

Table 2: Nowcast Evaluation Results

Nowcast Periods	Evaluation pre-pandemic				Evaluation incl. pandemic period			
	Average	6	13	18	Average	6	13	18
	RMSFE				RMSFE			
AR(2) benchmark (abs. RMSFE)	0.42	0.42	0.42	0.42	11.45	11.4	11.47	11.48
T-SV-t BMIDAS, GIGG w/ Spars. (rel. RMSFE)	0.66***	0.68**	0.60*	0.51**	0.21***	0.27	0.11	0.10
<i>Alternatives to T-SV-t (all BMIDAS, GIGG w/ Spars)</i>								
T-SV	0.75***	0.81*	0.72	0.44**	0.21***	0.27	0.17	0.08
T, Constant variance	0.76***	0.84	0.68	0.47**	0.21***	0.28	0.16	0.07
No T, SV-t	0.78***	0.90	0.70	0.46**	0.21***	0.31	0.10	0.08
No T, SV	0.81***	0.91	0.77	0.46**	0.22***	0.31	0.15	0.08
No T, Const. var.	1.03***	1.01	1.00*	0.67	0.22***	0.27	0.19	0.09
<i>Alternatives priors on MIDAS coefficients (all with T-SV-t)</i>								
GIGG w/out Spars.	0.69***	0.71*	0.62	0.49*	0.21***	0.32	0.10	0.09
Horseshoe (HS)	0.81***	0.87	0.73	0.74	0.28***	0.28	0.23	0.24
Spike and Slab (SS)	0.76***	0.74**	0.72*	0.78	0.32***	0.38	0.25	0.25
<i>Alternatives to BMIDAS (all with T-SV-t, GIGG w/ Spars)</i>								
U-BMIDAS	0.71***	0.74**	0.65**	0.60**	0.24***	0.29	0.18	0.18
MF-DFM	0.67***	0.75	0.67**	0.63**	0.32***	0.33	0.32	0.34
Combination univar. MIDAS	0.68***	0.68**	0.67**	0.67**	0.36***	0.37	0.34	0.33
	CRPS				CRPS			
AR(2) benchmark (abs. CRPS)	0.23	0.23	0.22	0.22	2.76	2.83	2.72	2.73
T-SV-t BMIDAS, GIGG w/ Spars. (rel. CRPS)	0.73***	0.76	0.68*	0.60*	0.26***	0.36	0.16	0.13
<i>Alternatives to T-SV-t (all BMIDAS, GIGG w/ Spars)</i>								
T-SV	0.75***	0.78	0.74*	0.48*	0.24***	0.31	0.19	0.11
T, Const. var.	0.84***	0.89	0.78*	0.63*	0.25***	0.32	0.19	0.12
No T, SV-t	0.82***	0.91	0.76*	0.53*	0.25***	0.36	0.15	0.11
No T, SV	0.83***	0.91	0.78*	0.50*	0.27***	0.37	0.18	0.11
No T, Const. var.	1.44***	1.31	1.60*	0.82*	0.31***	0.34	0.29	0.15
<i>Alternatives priors on MIDAS coefficients (all with T-SV-t)</i>								
GIGG w/out Spars.	0.77***	0.79	0.73*	0.61*	0.26***	0.36	0.16	0.13
Horseshoe (HS)	0.75***	0.79	0.67*	0.70	0.32***	0.31	0.26	0.26
Spike and Slab (SS)	0.75***	0.74	0.67*	0.73*	0.37***	0.43	0.28	0.27
<i>Alternatives to BMIDAS (all with T-SV-t, GIGG w/ Spars)</i>								
U-BMIDAS	0.79***	0.79	0.76*	0.75*	0.28***	0.33	0.22	0.20
MF-DFM	0.63***	0.60**	0.64**	0.60**	0.36***	0.36	0.36	0.36
Combination univar. MIDAS	0.68***	0.66**	0.68*	0.66**	0.41***	0.43	0.39	0.37

Notes: The table shows the average RMSFE and CRPS values for the AR model in the first row of each panel across all 20 nowcast periods (“Average”), and for selected nowcast periods (6,13,18). RMSFE and CRPS values for the other models are in relative terms to the AR model and stars indicate significance as per the Diebold-Mariano test Diebold et al. (1998) (* = 10% significance, ** = 5% significance, *** = 1% significance). T stands for time-varying trend, SV-t stands for stochastic volatility with t-distributed errors.

When including the Covid-19 pandemic (right panels), nowcast errors are higher for all models, particularly early on in the data release cycle, and differences across model variants are relatively smaller. Nowcast performance of all models clearly improves with the release of “hard” indicators for the first months of the reference quarter (nowcast period 13, i.e. 65 days prior to GDP), but more so for models that feature stochastic volatility and t-distributed errors. Interestingly, adding the time-varying trend makes less of a difference in the full evaluation

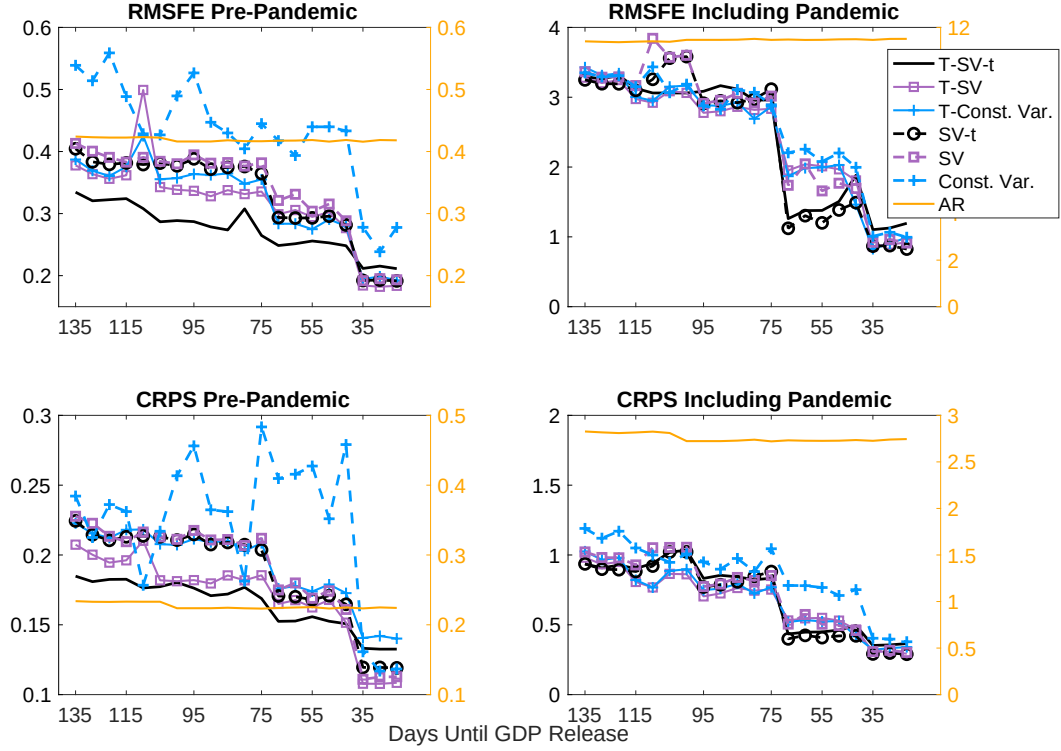


Figure 2: Nowcast performance against alternatives to T-SV-t component.

Notes: Absolute root mean square forecast errors (RMSFE) and continuous rank probability scores (CRPS) over nowcast periods (days ahead of GDP release). Left y-axis: baseline T-SV-t (black solid), T-SV (purple, square marker), T-Const Var (blue, plus), SV-t (dashed black, circle), SV (dashed purple, square), Const Var (dashed blue, plus). Right y-axis: AR(2) (yellow). All models use priors as outlined in section 2.2.

sample, and the full model loses out somewhat shortly before GDP release. The simple model without trend and with constant variance fares comparatively well early on in the data release cycle, but then loses out against the other models, particularly in terms of density nowcasts, its relative performance improves compared to pre-pandemic. Overall, this suggests that the most important time-varying feature when including the pandemic is the account for outliers, since it helps the model to identify the Covid-19 pandemic related downturn as temporary. On the other hand, the time-varying trend is less helpful during the pandemic, as models might over-fit the large shock into the trend if outliers are not accounted for.

To understand these results further, we look at the posterior trend and volatility estimates from the T-SV-t BMIDAS with GIGG prior. Figure 3 shows the posterior estimates of the cyclical and trend components (shown separately for pre-pandemic period and Covid-19 period for readability), and the stochastic volatility components of GDP growth and trend (lower panels). The trend-cycle decomposition is intuitive. The cyclical component captures high frequency movements in GDP growth and tracks actual GDP growth (black dashed-dotted lines) well, including over the Covid-19 pandemic, where the cyclical component captures the

bulk of the 20% drop in GDP growth and most of the recovery. On the other hand, the trend captures low frequency changes in UK GDP growth, with a gradual slowdown in trend growth since the early 2000s and a temporary trend decrease during the Great Financial Crisis. Throughout the pandemic, the trend remains largely unchanged. Further, the t-distributed volatility estimate of the observation equation shows a sharp and strong increase during the pandemic, by far exceeding the increase observed during the GFC. Hence, the model interprets the extreme movements in GDP growth as transitory in nature and related to a sharp rise in variance in GDP growth rather than in the long-run trend.

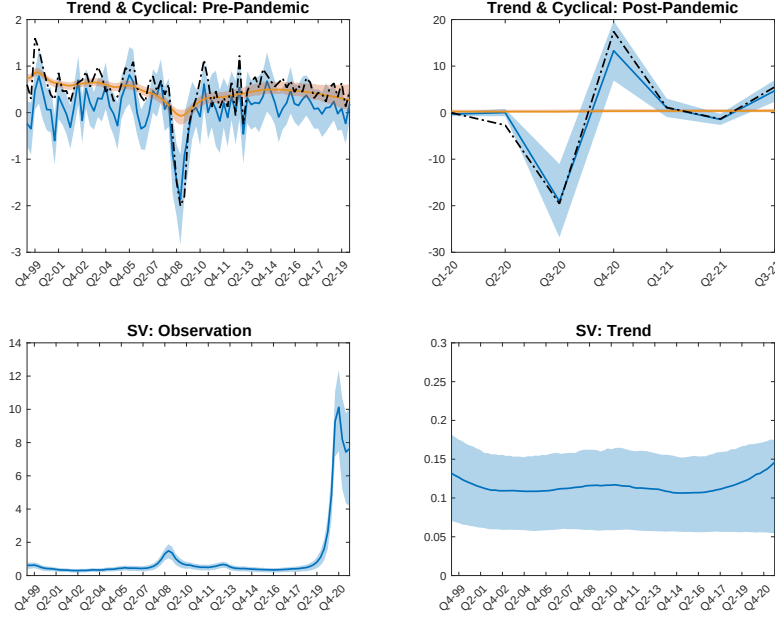


Figure 3: Posterior estimates for trend, cyclical component and stochastic volatilities.

Notes: Full sample, final nowcast period, from T-SV-t BMIDAS, with GIGG prior. Orange: posterior means for trend in GDP growth. Blue: posterior medians of the cyclical component (upper panel) and stochastic volatilities (lower panel). Black dashed line: UK GDP growth. Shaded areas: 95% credible intervals.

Separately identifying slow-moving trends from cyclical fluctuations in semi-structural models is generally challenging and can be sensitive to model choice, sample period and specification. The T-SV-t-BMIDAS is not free from this caveat, and therefore the economic interpretation that we can assign to the trend and cycle estimates can only be suggestive. This said, these flexible model features, in particularly stochastic volatility with t-distributed errors, appear to help achieving a more meaningful and stable trend-cycle identification across nowcast periods. Figure 4 shows the posterior trend and cyclical component from the baseline model compared to the alternative models, over the estimation period until 2019Q4 based on the information set at the first nowcast period, and the 18th nowcast period. The trend and cycle posterior estimates from the Trend-SV-t model remain similar across both nowcast periods, and reflect a

slow-moving trend and a cyclical component that captures fluctuations around a positive mean. By contrast, the model with SV but without t-distributed errors identifies a smooth cycle but volatile trend at the early nowcast period, which then shifts around at nowcast period 18. The model with constant variance identifies almost constant trends for both nowcast periods, but with huge uncertainty around them. For nowcast period 18, this model suggest a strongly negative almost fixed trend or intercept in GDP growth and cyclical fluctuations around strongly positive growth. These differences in the trend-cycle identification might explain the relatively weaker nowcast performance of these alternative models, and they suggest the model with SV and t-distributed errors can be preferable also in tranquil times prior to the pandemic.¹⁷

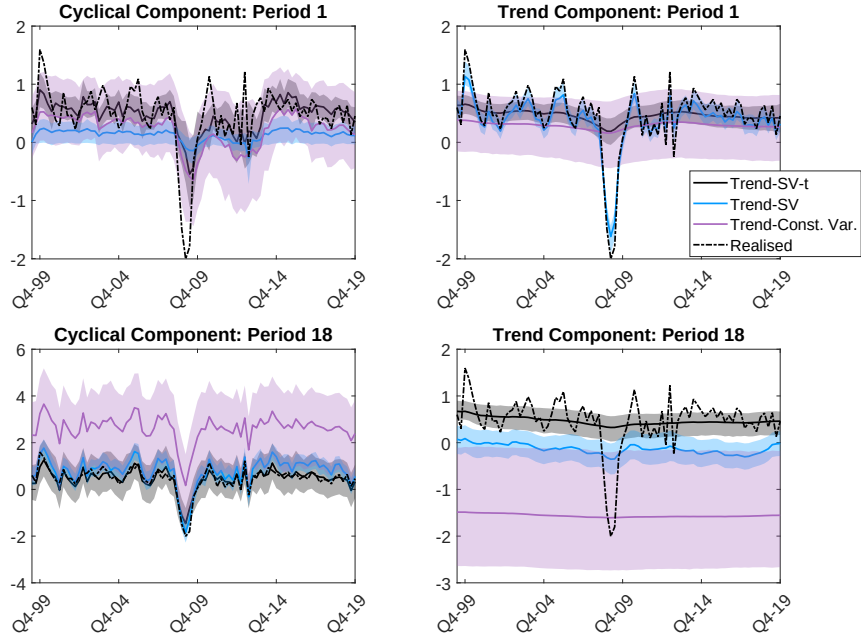


Figure 4: Posterior trend and cycle estimates across specifications.

Notes: Posterior means of trend estimated up until 2019Q4 from the T-SV-t BMIDAS model (black), the T-SV (blue), and T-Const.Var (purple). Estimation at the 1st and 18th nowcast periods. All models estimated with GIGG prior on the MIDAS coefficients.

4.2 Role of group-shrinkage prior with ex-post sparsification

Next, we look at the role of using the GIGG prior and the sparsification step for regularising the MIDAS regression, by comparing the nowcast results to similar models using alternative priors proposed in the literature. The alternative priors on MIDAS coefficients are

¹⁷ The fat-tailedness of error distributions in the stochastic volatility process proves to be an inherent model feature, as suggested by the posterior distribution of the degrees of freedom parameter showing large mass around small values, see Figure B1 in the appendix. Section B.4 of the appendix also show the trend-cycle identification against alternative priors for the time-varying components and MIDAS coefficients, suggesting a relatively more robust identification for the proposed model.

- *GIGG prior without sparsification step.*
- *Horseshoe prior (HS)*: model 1-3 with the horseshoe prior, thus another flexible shrinkage prior, but without group-shrinkage. We derive inclusion probabilities on a group-level using our group-sparsification algorithm from section 2.4.
- *Spike-and-slab prior (SSVS)*: This prior follows George and McCulloch (1993) with a uniform prior on lag-level inclusion probability. Selection of the prior is at the lag-level. We derive inclusion probabilities on a group-level using our group-sparsification algorithm.
- *Adaptive group-lasso with spike and slab prior (GAL)*, for Almon lag transformed data. This prior follows Mogliani and Simoni (2021). It is a group-wise spike-and-slab prior, where the slab distribution follows a multivariate Laplace distribution. Hyper-parameters that govern group-inclusion probability are found via an EM algorithm.

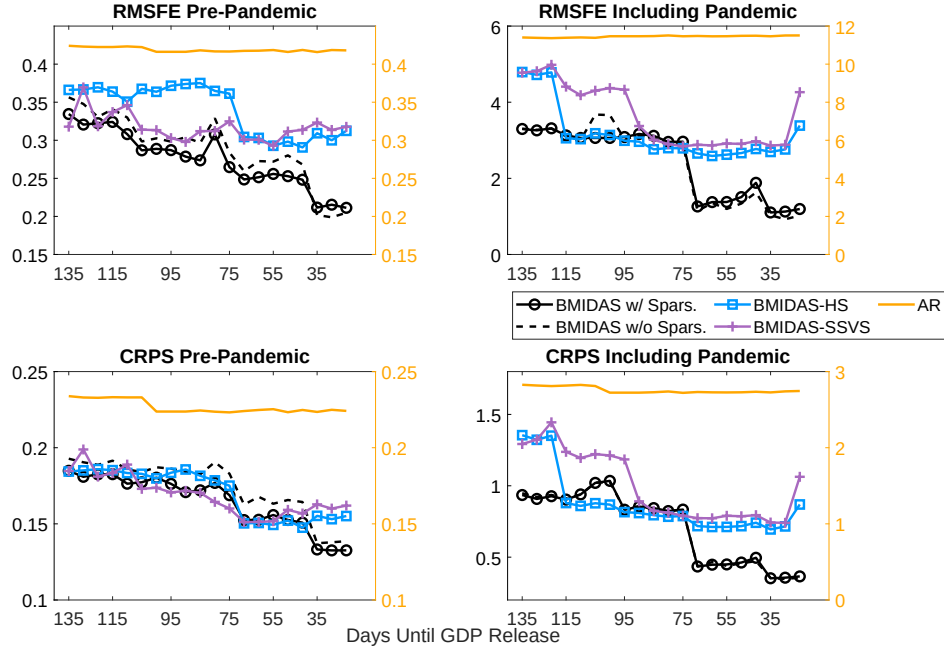


Figure 5: Nowcast performance across alternative priors, T-SV-t-BMIDAS.

Notes: Absolute RMSFE and CRPS for T-SV-t-BMIDAS models with GIGG prior and ex-post sparsification (black solid, circle marker), GIGG without sparsification (black dashed), Horseshoe prior HS (blue, square), Spike-and-slab prior SS (purple, plus). Right y-axis: AR(2) (yellow). All with time-varying trend, SV, t-distributed errors. Models with the GAL prior are only run without time-varying coefficients since the full model with that prior becomes computationally heavy.

First, we compare these prior specifications across models including the time-varying trend with stochastic volatility and t-distributed errors. Figure 5 shows the point and density nowcast evaluation, and Figure 6 visualises the nowcasts over time for the first and 13th nowcast periods. Exerting group-level shrinkage and taking into account the high-frequency correlation structure via the GIGG prior has preferable performance compared to the horseshoe prior and SSVS prior.

Prior to the pandemic, the model with GIGG prior has lower point nowcast errors throughout and almost always better density fit. The GIGG prior nowcasts are somewhat less volatile and show smaller uncertainty bands compared to the alternative priors. When including the pandemic quarters, it is strongest for the first nowcast periods and again starting from 75 days ahead of GDP publication (nowcast period 13 when “hard” indicators are released). All models initially miss the large unprecedented trough. However, in the first nowcast period, the proposed model with GIGG prior is the only model to indicate the large rebound in activity for Q3-2020. And once real activity indicators for the first reference month are released in nowcast period 13, the model with GIGG prior captures the downward adjustment for Q2-2020 most closely, but it also reflects the highly uncertain environment via the largest uncertainty bands. Nowcast uncertainty then decreases substantially after Q3-2020 for later nowcast periods, albeit it remains a bit larger than for the other models. Hence, with more “hard” macroeconomic information, the model indicates a return toward reduced uncertainty.

Finally, is it the addition of time-varying coefficients or the reliance on group-shrinkage that makes most of the difference for nowcast performance? Or does the group-shrinkage prior work particularly well when combined with the time-varying coefficients? Figure 7 shows results for specifications without trend and constant variance and for alternative priors on the MIDAS coefficients. Here, we also add a specification using the GAL prior, as proposed by Mogliani and Simoni (2021). We see that without time-varying components, the GIGG prior performs rather similarly to the horseshoe and GAL priors, and for density forecasts and when including the pandemic it loses some of its gains seen from Figure 5. The SSVS prior has a weaker and volatile performance without time-variation included, particularly prior to the pandemic. Including the pandemic, the models without time-varying components both with the GIGG prior or the GAL prior achieve some improvement around nowcast period 13 when “hard” data are released, but much less so than the model with time-varying components and GIGG. At nowcast period 18, the models with GIGG prior with and without time-variation perform equally well.

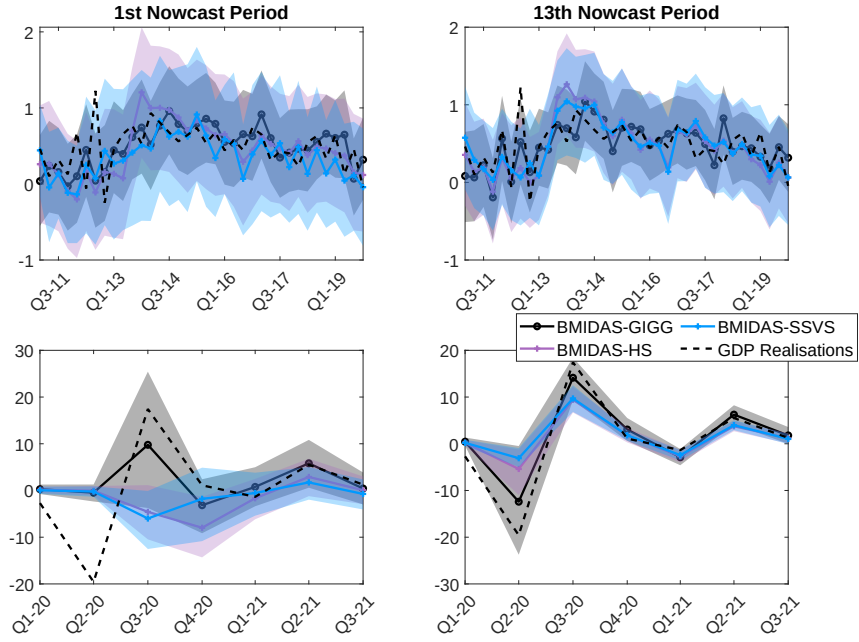


Figure 6: Posterior mean and density nowcasts, for alternatives priors on MIDAS component. Notes: Nowcasts over period 2011Q1 to 2019Q4 (x-axis) for the 1st and 13th nowcast periods, and over the pandemic quarters (lower panel). Models are as in Figure 5. Shaded areas show 95% credible intervals. Black dashed lines show quarterly GDP growth realisations.

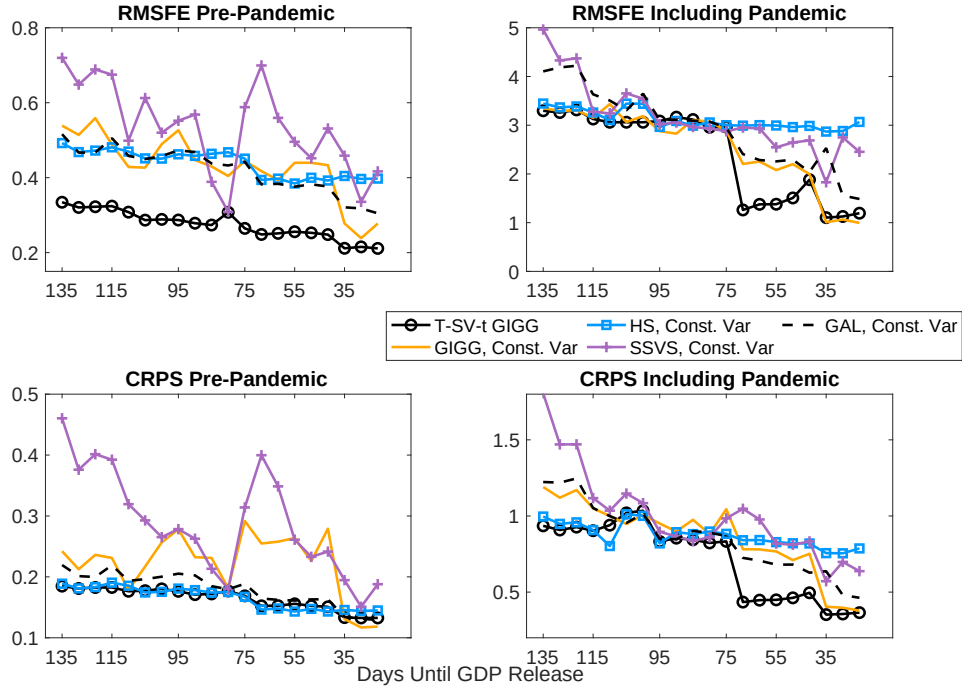


Figure 7: Alternative priors for BMIDAS without trend and constant variance. Notes: RMSFE and CRPS for baseline T-SV-t BMIDAS with GIGG (solid black line, circle marker), and for BMIDAS models without trend and constant volatility and with different priors on MIDAS coefficients: GIGG (yellow), horseshoe (blue, square marker), SSVS (purple, plus marker), GAL (dashed black).

Overall, the results suggests that the combination of time-varying components and group-shrinkage prior improves nowcast performance over both evaluation periods. As we illustrate in the next section, the GIGG prior shrinks the information set towards a sparse selection of indicators, while still drawing on other indicators to a much lesser extent. And in a model with time-varying coefficients the prior is also able to shift among the most meaningful indicators over the data release cycle and over time, rather than relying on constant signals.

5 Interpreting signals from groups of indicators over time

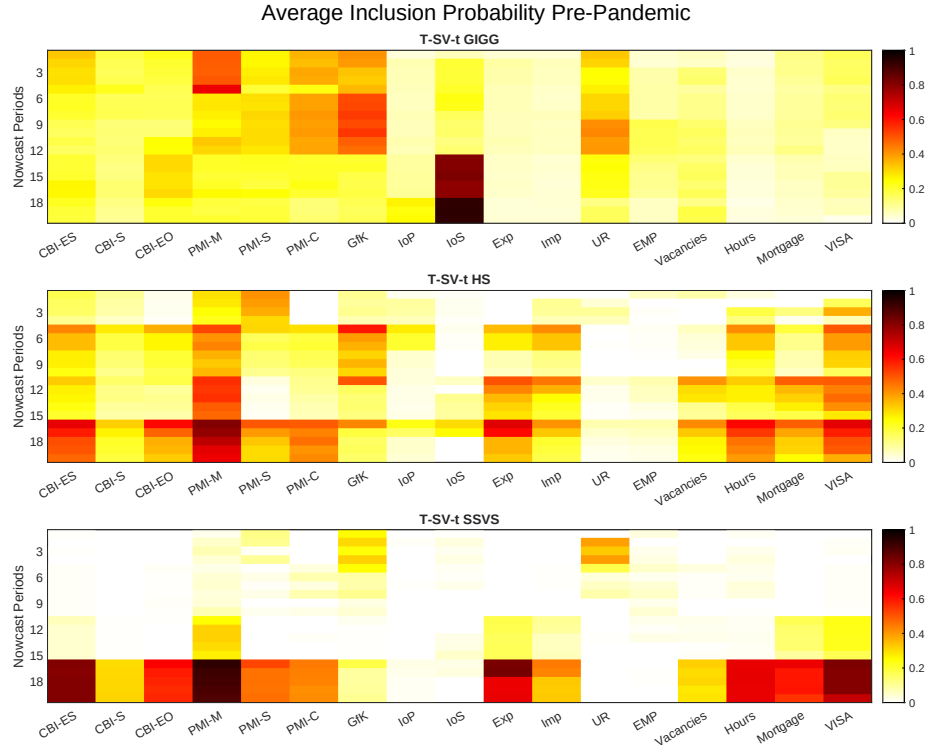
The sparsification algorithm (9) on the posterior achieves variable selection that allows us to single out which indicators have the largest impact on the predictions of the model over time and over the data release cycle. Figure 9 presents heatmaps for inclusion probabilities for groups of high-frequency lags belonging to each indicator (x-axis) and nowcast period (y-axis), for the pre-pandemic evaluation sample (panel a) and the full evaluation sample (panel b).¹⁸ The upper sub-plots in both panels show the T-SV-t BMIDAS with GIGG prior, and the lower sub-plots results from the same model but using the horseshoe prior without group-shrinkage, and the spike-and-slab SSVS prior.

The model with GIGG prior mainly relies on sparse set of indicators at a time, as indicated by only a few red-shaded areas. However, the inclusion probabilities of all other indicators are non-zero as well, as indicated by the yellow colours of different intensity. Thus, information is exploited in line with the idea of “illusion of sparsity” (Giannone et al., 2021), where a broad set of indicators are potentially useful for forecasts, but many signals will be small. Interestingly, the model shifts between groups of indicators over the data release cycle (i.e. moving down the y-axis). Early on, the model exploits a few survey indicators. For the pre-pandemic evaluation sample, these are manufacturing and construction PMIs for very early nowcasts, and then GfK consumer confidence once its observations for the first month of the reference quarter are released. This agrees with earlier findings that survey indicators provide the main early signals for quarterly GDP (Bańbura et al., 2013; Anesti et al., 2017). Later on, labour market data also plays a role when released in period 9.

For nowcast period 13, once “hard” economic information get published, we had seen a clear improvement in nowcast performance of the proposed model in Figure 5. The inclusion probabilities reveal that from this period on, the model relies almost exclusively on the index of services, an important indicator for the service-oriented UK economy. Instead, the model with horseshoe prior shows dense and a quite diffuse inclusion pattern across indicators and over nowcast periods, and for both evaluation samples. It draws on signals from surveys, real,

¹⁸ Inclusion probabilities are stable over time apart from a pandemic-induced shift, so we focus on averages over sub-samples. See Table B2 in the appendix for the exact inclusion probabilities and standard deviations.

a) Pre-Pandemic (2011Q1-2019Q4)



b) Including the Pandemic (2011Q1-2021Q2)

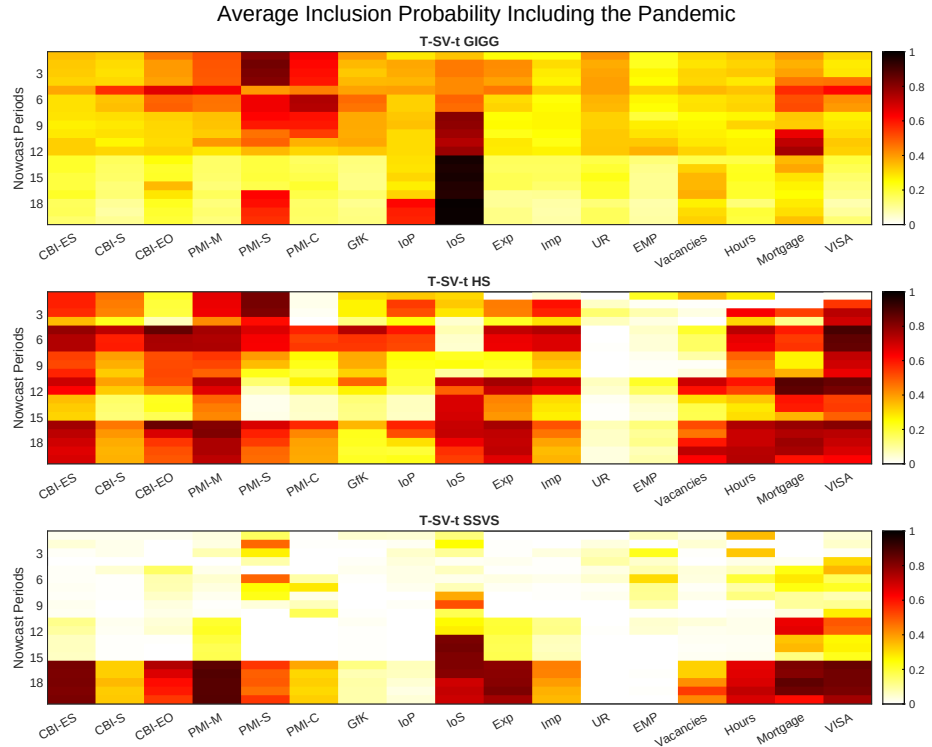


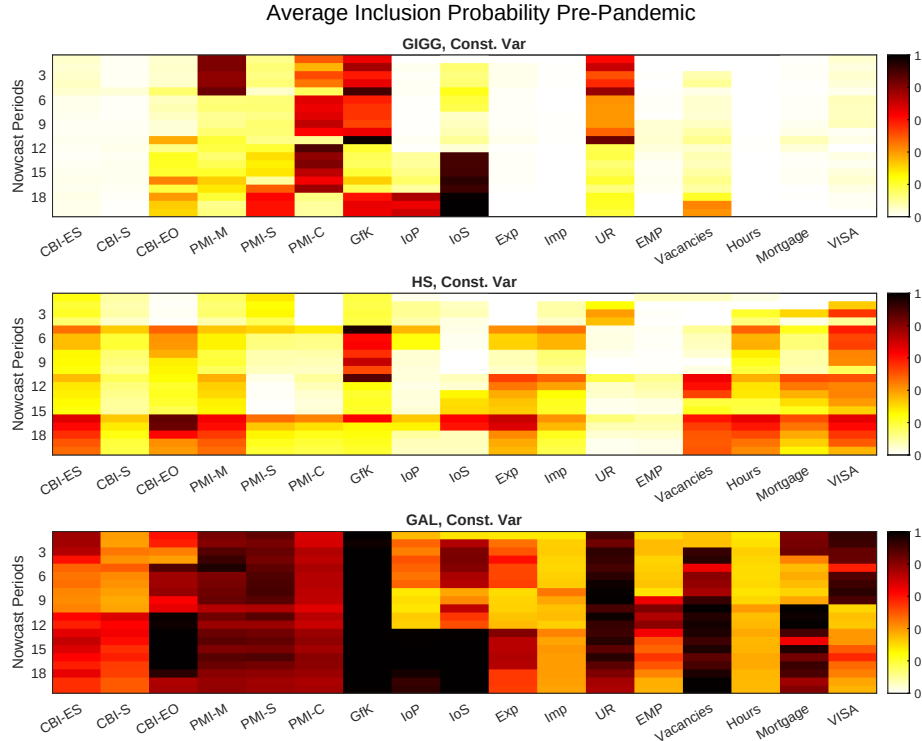
Figure 8: Inclusion probabilities across prior choices, BMIDAS models with trend and SV-t. Notes: Darker colour indicates higher cumulative posterior inclusion probability for groups of monthly lags of each indicator (x-axis) over nowcast periods (y-axis), from T-SV-t BMIDAS model. Panel a) shows average incl. prob. for evaluation until 2019Q4, panel b) until 2021Q3. Sub-plots show the GIGG prior and horseshoe prior HS, respectively.

labour and personal finance indicators, but switches between these indicators over the data release cycle, such as no clear pattern emerges. The SSVS prior, on the other hand, reads little information in early stage of the data release cycle, and only reads intense but broad based signal close to actual GDP release. It therefore seems less able to exploit timely information early on. Overall, by switching from one set of indicators to another, the proposed model with GIGG prior seems to exploit the incoming information more efficiently compared to other models that rely on a more “dense” reading of information.

With the pandemic included in the evaluation, the higher volatility leads to more intense colours as more signals are being processed, independently of the prior. For the model with GIGG prior, the clear pattern remains of exploiting different signals over the data release cycle, while still incorporating all other indicators to a much lesser extent. Survey indicators remain most important for early nowcast periods, these are now rather the service and construction PMIs, less so manufacturing. Over nowcast periods 5 to 12, signals from mortgage lending are relevant too, and the index of services turns important, even before its observations for the current quarter get released. Instead, little focus is put on the GFK and labour market data. Thus, the model is able to capture that during the pandemic the economic lockdowns affected mainly the service sector and initially the housing and construction sectors, whereas consumer confidence and manufacturing were less affected and labour market data were distorted by the furlough scheme. In nowcast period 13, there is again a shift towards a strong focus on the index of services, and the inclusion probabilities of most other indicators drop. However, now the model also continues relying on vacancies and mortgage lending, and in nowcast period 18 additional signals stem from the index of production, and again from the service PMI. The horseshoe prior now also relies less on labour market series compared to the pre-pandemic period, and somewhat more on mortgages and VISA consumer spending, but it does not switch away from manufacturing PMIs and other survey indicators. Given that the Covid-19 shock affected specific sectors more than others, a dense solution can represent a disadvantage, as suggested by the weaker nowcast performance of the model with horseshoe prior.

Finally, what is the role of time-varying components for the observed inclusion patterns? Figure 9 shows a similar comparison of inclusion probabilities for the models without time-varying trend and with constant variance, with the GIGG prior being compared against the horseshoe and the GAL priors. Overall, the inclusion patterns remain similar. Two notable differences for the model with GIGG prior stand out compared to the case without time-varying components. First, in the model with constant coefficients the GIGG prior is more selective: it selects sub-groups of indicators and the inclusion probabilities of other indicators are close to zero, particularly prior to the pandemic. Thus, while the aggressive group shrinkage helps the model rely on a sub-set of most informative indicators more than on others, the account for the time-varying trend nudges the model to still read weaker signals from all other indicators as well,

a) Pre-Pandemic (2011Q1-2019Q4)



b) Including the Pandemic (2011Q1-2021Q2)

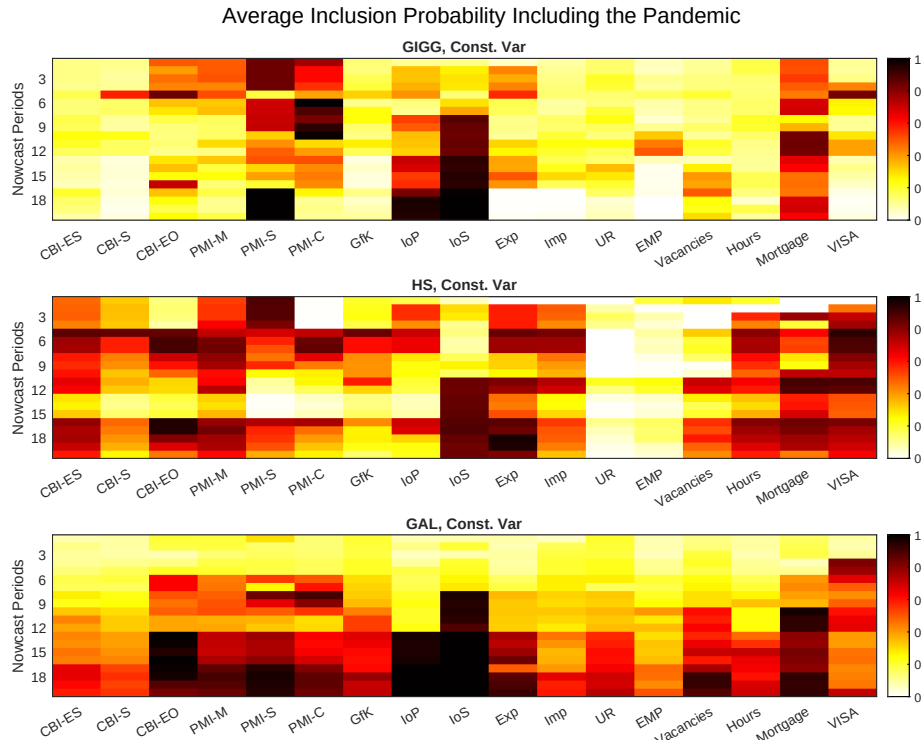


Figure 9: Inclusion probabilities, BMIDAS models without trend and constant variance. Notes: Darker colour indicates higher cumulative posterior inclusion probability for groups of monthly lags of each indicator (x-axis) over nowcast periods (y-axis), from BMIDAS models with constant variance and no trend. Panel a) shows average incl. prob. for evaluation until 2019Q4, panel b) until 2021Q3. Sub-plots show, respectively, the GIGG prior, horseshoe prior, and group adaptive lasso with spike and slab (GAL) prior as in Mogliani and Simoni (2021).

in line with the “illusion of sparsity” argument. Second, the constant coefficients model does not show a pronounced shift in nowcast period 13 towards a sparser specification; while some probabilities change, the intensity of most colours remains the same. Thus, the time-varying components likely help the model to shift towards signals from “hard” data once these become available. This is not the case for priors without group-shrinkage such as the horseshoe, with or without time-varying components. Overall, this illustrates how efficient group-shrinkage, together with the inclusion of time-varying coefficients, can help select the most meaningful indicators over time and the data release cycle and jointly improve nowcast performance.

6 Conclusion

We have proposed a new Bayesian MIDAS framework, the T-SV-t-BMIDAS model combined with a flexible group-shrinkage prior that accounts for the correlation within groups of indicators, in the MIDAS case within groups of high-frequency lags. Group-shrinkage and sparsification are achieved in separation. The former serves to regularise the MIDAS coefficients, whereas the latter achieves interpretability and communication of results via a sparsification step on the posterior motivated by Bayesian decision theory.

In an application of the model to nowcasting UK GDP growth, our model is able to capture sharp changes in GDP growth as seen during the Covid-19 pandemic in a relatively timely manner, and works well also in more tranquil times, particularly at a later stage of the data release cycle when “hard” economic indicators are released. The aggressive group-shrinkage helps selecting the most relevant indicators and works particularly well when combined with time-varying components. When accounting for a time-varying trend in GDP growth, the prior achieves group-sparsity by selecting the most relevant indicators, but it also reads a wider range of much smaller signals from all other indicators, in line with the “illusion of sparsity” argument. Allowing for t-distributed errors proves particularly relevant once the Covid-19 pandemic is included. It helps the prior to shift between groups of indicators over time and over the data release cycle, which can be important in presence of a large heterogeneous shock.

The proposed model is flexible, combining features and nesting various models as special cases. The derived inclusion probabilities make results more interpretable and easier to communicate over time and over the data release cycle. All this makes an attractive tool not only for nowcasting GDP growth, but also for a wider range of forecasting applications. The group-shrinkage prior has potential to efficiently regularise information towards a group-sparse set of signals for a range of other data sets where groupings might be present, such as across sectors or countries. We leave the exploration of such alternative applications for future research.

References

- Almon, S. (1965). The distributed lag between capital appropriations and expenditures. *Econometrica: Journal of the Econometric Society*, 178–196.
- Andreou, E., E. Ghysels, and A. Kourtellis (2010). Regression models with mixed sampling frequencies. *Journal of Econometrics* 158(2), 246–261.
- Anesti, N., A. B. Galvão, and S. Miranda-Agrippino (2018). Uncertain Kingdom: nowcasting GDP and its revisions. Bank of England Working Paper No 764.
- Anesti, N., S. Hayes, A. Moreira, and J. Tasker (2017). Peering into the present: the Bank’s approach to GDP nowcasting. *Bank of England Quarterly Bulletin*, Q2.
- Angelini, E., G. Camba-Mendez, D. Giannone, L. Reichlin, and G. Rünstler (2011). Short-term forecasts of euro area gdp growth.
- Antolin-Diaz, J., T. Drechsel, and I. Petrella (2017). Tracking the slowdown in long-run GDP growth. *Review of Economics and Statistics* 99(2), 343–356.
- Antolin-Diaz, J., T. Drechsel, and I. Petrella (2021). Advances in nowcasting economic activity: Secular trends, large shocks and new data. CEPR Discussion Paper No. DP15926.
- Armagan, A., D. B. Dunson, and J. Lee (2013). Generalized double Pareto shrinkage. *Statistica Sinica* 23(1), 119.
- Babii, A., R. T. Ball, E. Ghysels, and J. Striaukas (2022, July). Machine learning panel data regressions with heavy-tailed dependent data: Theory and application. *Journal of Econometrics*.
- Babii, A., E. Ghysels, and J. Striaukas (2022, March). Machine learning time series regressions with an application to nowcasting. *Journal of Business & Economic Statistics* 40(3), 1094–1106.
- Bai, J., E. Ghysels, and J. H. Wright (2013). State space models and midas regressions. *Econometric Reviews* 32(7), 779–813.
- Bañbura, M., D. Giannone, M. Modugno, and L. Reichlin (2013). Now-casting and the real-time data flow. In *Handbook of economic forecasting*, Volume 2, pp. 195–237. Elsevier.
- Bañbura, M. and M. Modugno (2014). Maximum likelihood estimation of factor models on datasets with arbitrary pattern of missing data. *Journal of Applied Econometrics* 29(1), 133–160.
- Bank of England, . (2020). How are we monitoring the economy during the Covid-19 pandemic? Bank Overground, <https://www.bankofengland.co.uk/bank-overground/2020/how-are-we-monitoring-the-economy-during-the-covid-19-pandemic>.
- Barbieri, M. M. and J. O. Berger (2004). Optimal predictive model selection. *The annals of statistics* 32(3), 870–897.
- Barbieri, M. M., J. O. Berger, E. I. George, and V. Ročková (2021). The median probability model and correlated variables. *Bayesian Analysis* 16(4), 1085–1112.
- Baumeister, C., D. Leiva-León, and E. R. Sims (2021). Tracking weekly state-level economic conditions. Technical report, National Bureau of Economic Research.
- Berger, T., G. Everaert, and H. Vierke (2016). Testing for time variation in an unobserved components model for the us economy. *Journal of Economic Dynamics and Control* 69, 179–208.
- Bhattacharya, A., A. Chakraborty, and B. K. Mallick (2016). Fast sampling with gaussian scale mixture priors in high-dimensional regression. *Biometrika* 103(4), 985–991.
- Boss, J., J. Datta, X. Wang, S. K. Park, J. Kang, and B. Mukherjee (2021). Group Inverse-

- Gamma Gamma Shrinkage for Sparse Regression with Block-Correlated Predictors. *arXiv preprint arXiv:2102.10670*.
- Breheny, P. and J. Huang (2015). Group descent algorithms for nonconvex penalized linear and logistic regression models with grouped predictors. *Statistics and Computing* 25(2), 173–187.
- Carriero, A., T. E. Clark, and M. Marcellino (2015). Realtime nowcasting with a Bayesian mixed frequency model with stochastic volatility. *Journal of the Royal Statistical Society. Series A, (Statistics in Society)* 178(4), 837.
- Carriero, A., T. E. Clark, and M. G. Marcellino (2016). Large vector autoregressions with stochastic volatility and flexible priors.
- Carriero, A., T. E. Clark, M. G. Marcellino, and E. Mertens (2021). Addressing COVID-19 outliers in BVARs with stochastic volatility. CEPR Discussion Paper No. DP15964.
- Carriero, A., A. B. Galvão, and G. Kapetanios (2019). A comprehensive evaluation of macroeconomic forecasting methods. *International Journal of Forecasting* 35(4), 1226–1239.
- Carter, C. K. and R. Kohn (1994). On Gibbs sampling for state space models. *Biometrika* 81(3), 541–553.
- Carvalho, C. M., N. G. Polson, and J. G. Scott (2010). The horseshoe estimator for sparse signals. *Biometrika* 97(2), 465–480.
- Casella, G., M. Ghosh, J. Gill, and M. Kyung (2010). Penalized regression, standard errors, and Bayesian lassos. *Bayesian Analysis* 5(2), 369–411.
- Chakraborty, A., A. Bhattacharya, and B. K. Mallick (2020). Bayesian sparse multiple regression for simultaneous rank reduction and variable selection. *Biometrika* 107(1), 205–221.
- Chan, J. C. (2017). Notes on Bayesian macroeconometrics. Manuscript available at <http://joshuachan.org>.
- Chan, J. C. and A. L. Grant (2016). Modeling energy price dynamics: Garch versus stochastic volatility. *Energy Economics* 54, 182–189.
- Chan, J. C. and I. Jeliazkov (2009). Efficient simulation and integrated likelihood estimation in state space models. *International Journal of Mathematical Modelling and Numerical Optimisation* 1(1-2), 101–120.
- Chiu, C.-W. J., H. Mumtaz, and G. Pinter (2017). Forecasting with var models: Fat tails and stochastic volatility. *International Journal of Forecasting* 33(4), 1124–1143.
- Clark, T. E. (2011). Real-time density forecasts from bayesian vector autoregressions with stochastic volatility. *Journal of Business & Economic Statistics* 29(3), 327–341.
- Clark, T. E. and F. Ravazzolo (2015). Macroeconomic forecasting performance under alternative specifications of time-varying volatility. *Journal of Applied Econometrics* 30(4), 551–575.
- Cogley, T., S. Morozov, and T. J. Sargent (2005). Bayesian fan charts for UK inflation: Forecasting and sources of uncertainty in an evolving monetary system. *Journal of Economic Dynamics and Control* 29(11), 1893–1925.
- D’Agostino, A., L. Gambetti, and D. Giannone (2013). Macroeconomic forecasting and structural change. *Journal of applied econometrics* 28(1), 82–101.
- Devroye, L. (2014). Random variate generation for the generalized inverse gaussian distribution. *Statistics and Computing* 24(2), 239–246.
- Diebold, F. X., T. A. Gunther, and A. S. Tay (1998). Evaluating density forecasts with applications to financial risk management. *International Economic Review* 39(4), 863–883.
- Ferrara, L., M. Mogliani, and J.-G. Sahuc (2022). High-frequency monitoring of growth at risk. *International Journal of Forecasting* 38(2), 582–595.
- Foroni, C. and M. Marcellino (2014). A comparison of mixed frequency approaches for now-

- casting euro area macroeconomic aggregates. *International Journal of Forecasting* 30(3), 554–568.
- Foroni, C., M. Marcellino, and C. Schumacher (2015). Unrestricted mixed data sampling (MIDAS): MIDAS regressions with unrestricted lag polynomials. *Journal of the Royal Statistical Society: Series A (Statistics in Society)* 178(1), 57–82.
- Friedman, J., T. Hastie, and R. Tibshirani (2010). Regularization paths for generalized linear models via coordinate descent. *Journal of Statistical Software* 33(1), 1.
- Frühwirth-Schnatter, S. and H. Wagner (2010). Stochastic model specification search for gaussian and partial non-gaussian state space models. *Journal of Econometrics* 154(1), 85–100.
- George, E. I. and R. E. McCulloch (1993). Variable selection via Gibbs sampling. *Journal of the American Statistical Association* 88(423), 881–889.
- Ghysels, E., V. Kvedaras, and V. Zemlys-Balevičius (2020). Mixed data sampling (midas) regression models. In *Handbook of Statistics*, Volume 42, pp. 117–153. Elsevier.
- Ghysels, E., P. Santa-Clara, and R. Valkanov (2004). The midas touch: Mixed data sampling regression models.
- Ghysels, E., A. Sinko, and R. Valkanov (2007). MIDAS regressions: Further results and new directions. *Econometric reviews* 26(1), 53–90.
- Giannone, D., M. Lenza, and G. E. Primiceri (2019). Priors for the long run. *Journal of the American Statistical Association* 114(526), 565–580.
- Giannone, D., M. Lenza, and G. E. Primiceri (2021). Economic predictions with big data: The illusion of sparsity. *Econometrica* 89(5), 2409–2437.
- Gneiting, T. and A. E. Raftery (2007). Strictly proper scoring rules, prediction, and estimation. *Journal of the American Statistical Association* 102(477), 359–378.
- Griffin, J. E. and P. J. Brown (2012). Structuring shrinkage: some correlated priors for regression. *Biometrika* 99(2), 481–487.
- Guérin, P. and M. Marcellino (2013). Markov-switching midas models. *Journal of Business & Economic Statistics* 31(1), 45–56.
- Hahn, P. R. and C. M. Carvalho (2015). Decoupling shrinkage and selection in bayesian linear models: a posterior summary perspective. *Journal of the American Statistical Association* 110(509), 435–448.
- Harvey, A. C., T. M. Trimbur, and H. K. Van Dijk (2007). Trends and cycles in economic time series: A bayesian approach. *Journal of Econometrics* 140(2), 618–649.
- Hauzenberger, N., F. Huber, and L. Onorante (2021). Combining shrinkage and sparsity in conjugate vector autoregressive models. *Journal of Applied Econometrics* 36(3), 304–327.
- Hörmann, W. and J. Leydold (2014). Generating generalized inverse Gaussian random variates. *Statistics and Computing* 24(4), 547–557.
- Huber, F., G. Koop, and L. Onorante (2021). Inducing sparsity and shrinkage in time-varying parameter models. *Journal of Business & Economic Statistics* 39(3), 669–683.
- Huber, F., G. Koop, L. Onorante, M. Pfarrhofer, and J. Schreiner (2020). Nowcasting in a pandemic using non-parametric mixed frequency VARs. *Journal of Econometrics*.
- Jacquier, E., N. G. Polson, and P. E. Rossi (2004). Bayesian analysis of stochastic volatility models with fat-tails and correlated errors. *Journal of Econometrics* 122(1), 185–212.
- Kapetanios, G., F. Papailias, et al. (2022). Real Time Indicators During the COVID-19 Pandemic Individual Predictors & Selection. Economic Statistics Centre of Excellence (ESCoE).
- Kim, C.-J. and C. R. Nelson (1999). Has the US economy become more stable? A Bayesian approach based on a Markov-switching model of the business cycle. *Review of Economics*

- and *Statistics* 81(4), 608–616.
- Kim, S., N. Shephard, and S. Chib (1998). Stochastic volatility: likelihood inference and comparison with ARCH models. *The Review of Economic Studies* 65(3), 361–393.
- Kohns, D. and A. Bhattacharjee (2022). Nowcasting growth using google trends data: A bayesian structural time series model. *International Journal of Forecasting*.
- Lenza, M. and G. E. Primiceri (2022). How to estimate a vector autoregression after March 2020. *Journal of Applied Econometrics*.
- Lindley, D. V. (1968). The choice of variables in multiple regression. *Journal of the Royal Statistical Society: Series B (Methodological)* 30(1), 31–53.
- Makalic, E. and D. F. Schmidt (2015). A simple sampler for the horseshoe estimator. *IEEE Signal Processing Letters* 23(1), 179–182.
- McCausland, W. J., S. Miller, and D. Pelletier (2011). Simulation smoothing for state-space models: A computational efficiency analysis. *Computational Statistics & Data Analysis* 55(1), 199–212.
- McConnell, M. M. and G. Perez-Quiros (2000). Output fluctuations in the United States: What has changed since the early 1980’s? *American Economic Review* 90(5), 1464–1476.
- Mogliani, M. and A. Simoni (2021). Bayesian MIDAS penalized regressions: estimation, selection, and prediction. *Journal of Econometrics* 222(1), 833–860.
- Ng, S. (2021). Modeling macroeconomic variations after COVID-19. NBER Working Paper No 29060.
- O’Hara, R. B. and M. J. Sillanpää (2009). A review of bayesian variable selection methods: what, how and which. *Bayesian analysis* 4(1), 85–117.
- Piironen, J., M. Paasiniemi, A. Vehtari, et al. (2020). Projective inference in high-dimensional problems: Prediction and feature selection. *Electronic Journal of Statistics* 14(1), 2155–2197.
- Piironen, J. and A. Vehtari (2017). Comparison of bayesian predictive methods for model selection. *Statistics and Computing* 27(3), 711–735.
- Piironen, J., A. Vehtari, et al. (2017). Sparsity information and regularization in the horseshoe and other shrinkage priors. *Electronic Journal of Statistics* 11(2), 5018–5051.
- Polson, N. G. and J. G. Scott (2010). Shrink globally, act locally: Sparse Bayesian regularization and prediction. *Bayesian Statistics* 9, 501–538.
- Ray, P. and A. Bhattacharya (2018). Signal adaptive variable selector for the horseshoe prior. *arXiv preprint arXiv:1810.09004*.
- Scruton, J., M. O’Donnell, and S. Dey-Chowdhury (2018). Introducing a new publication model for GDP. *Office for National Statistics*, 1–14.
- Simon, N. and R. Tibshirani (2012). Standardization and the group lasso penalty. *Statistica Sinica* 22(3), 983.
- Smith, R. G. and D. E. Giles (1976). *The Almon estimator: Methodology and users’ guide*. Reserve Bank of New Zealand.
- Stock, J. H. and M. W. Watson (2004). Combination forecasts of output growth in a seven-country data set. *Journal of Forecasting* 23(6), 405–430.
- Stock, J. H. and M. W. Watson (2007). Why has US inflation become harder to forecast? *Journal of Money, Credit and Banking* 39, 3–33.
- Stock, J. H. and M. W. Watson (2009). Forecasting in dynamic factor models subject to structural instability. *The Methodology and Practice of Econometrics. A Festschrift in Honour of David F. Hendry* 173, 205.
- Stock, J. H. and M. W. Watson (2016). Core inflation and trend inflation. *Review of Economics*

- and Statistics* 98(4), 770–784.
- Wang, H. and C. Leng (2008). A note on adaptive group lasso. *Computational statistics & data analysis* 52(12), 5277–5286.
- Woloszko, N. (2020). Tracking activity in real time with Google Trends. OECD Working Paper No 1634.
- Woody, S., C. M. Carvalho, and J. S. Murray (2021). Model interpretation through lower-dimensional posterior summarization. *Journal of Computational and Graphical Statistics* 30(1), 144–161.
- Xu, X. and M. Ghosh (2015). Bayesian variable selection and estimation for group lasso. *Bayesian Analysis* 10(4), 909–936.
- Xu, Z., D. F. Schmidt, E. Makalic, G. Qian, and J. L. Hopper (2016). Bayesian grouped horseshoe regression with application to additive models. In *Australasian Joint Conference on Artificial Intelligence*, pp. 229–240. Springer.
- Zou, H. (2006). The adaptive lasso and its oracle properties. *Journal of the American Statistical Association* 101(476), 1418–1429.

A Model Descriptions

A.1 Posteriors

Here we list the conditional posteriors derived using standard method from the literature (referenced where needed). These allow implimenting the Gibbs sampler in [reference to Gibbs sampler]

Posterior for θ

Consider the likelihood in 7 and define the stacked prior for the MIDAS coefficients as $p(\theta|\vartheta, \gamma^2, \varphi^2) = N(0, \vartheta\Gamma\varphi^2)$, where $\Gamma = \text{diag}(\gamma_1^2, \dots, \gamma_1^2, \dots, \gamma_K^2, \dots, \gamma_K^2)$ and $\varphi^2 = \text{diag}(\varphi_{1,1}^2, \dots, \varphi_{1,L-r}^2, \dots, \varphi_{K,1}^2, \dots, \varphi_{K,L-r}^2)$. Denote the prior covariance and the covariance of the target for simplicity as Λ_* and $\Lambda_{\lambda,h} = \Lambda_\lambda\Lambda_h$ respectively. Then, the posterior by standard manipulations is:

$$\begin{aligned} \theta|\mathbf{y}, \bullet &\sim N(\bar{\theta}, \bar{\Lambda}_*^{-1}) \\ \bar{\Lambda}_* &= (\mathbf{Z}'\Lambda_{\lambda,h}^{-1}\mathbf{Z} + \Lambda_*^{-1}), \quad \bar{\theta} = \bar{\Lambda}_*^{-1}(\mathbf{Z}'\Lambda_{\lambda,h}^{-1}\tilde{\mathbf{y}}), \end{aligned} \quad (13)$$

Posteriors of Hyper-parameters

The derivations of the conditional posteriors for $\vartheta, \gamma_k^2, \varphi_{kj}$ for $k = 1, \dots, K$ and $j = 1, \dots, L-r$ immediately follow from the presentation in Boss et al. (2021). Following Boss et al. (2021), we employ a mixture representation of the β' prior via an inverse-gamma distributed auxiliary variable ν_p . The conditional posteriors $\vartheta, \gamma_k^2, \varphi_{kj}^2, \nu_p$ are thus proportional to:

$$(\vartheta|\mathbf{y}, \bullet) \sim IG\left(\frac{\sum_{k=1}^K(L-r) + 1}{2}, \theta'\Lambda_p^{-1}\theta/2 + \frac{1}{\nu_p}\right) \quad (14)$$

$$(\gamma_k^{-2}|\mathbf{y}, \bullet) \sim GIG\left(\frac{L-r}{2} - a_k, \frac{1}{\vartheta^2} \sum_{j=1}^{L-r} \frac{\theta_{kj}^2}{\varphi_{kj}^2}\right) \quad (15)$$

$$(\varphi_{kj}|\mathbf{y}, \bullet) \sim IG\left(b_k + \frac{1}{2}, 1 + \frac{\theta_{kj}^2}{2\vartheta^2\gamma_k^2}\right) \quad (16)$$

$$(\nu_p|\mathbf{y}, \bullet) \sim IG\left(1, \frac{1}{\vartheta}\right), \quad (17)$$

where GIG refers to the generalised inverse Gaussian distribution (Hörmann and Leydold, 2014) which we generate from using the efficient algorithm of (Devroye, 2014).

Posterior of τ

To derive the joint posterior of $\boldsymbol{\tau}$, we make use of the methods proposed by Chan and Jeliazkov (2009); McCausland et al. (2011) as they enable sampling of all states $(\tau_1, \dots, \tau_T)$ simultaneously. Define $\mathbf{y}^* = \mathbf{y} - \mathbf{Z}\theta$ and $\Lambda_g = \text{diag}(e^{g_1}, \dots, e^{g_T})$. Let H be the first difference matrix, then:

$$\boldsymbol{\tau} | \mathbf{y}, \tau_0, \bullet \sim N(\hat{\boldsymbol{\tau}}, K_{\boldsymbol{\tau}}^{-1}), \quad (18)$$

where $K_{\boldsymbol{\tau}} = H' \Lambda_g^{-1} H + \Lambda_h^{-1} \Lambda_{\lambda}^{-1/2}$, $\hat{\boldsymbol{\tau}} = K_{\boldsymbol{\tau}}^{-1} (H \Lambda_g^{-1} H \tau_0 \mathbf{1}_T + \Lambda_h^{-1} \Lambda_{\lambda}^{-1/2} \mathbf{y}^*)$. And,

$$\tau_0 | \mathbf{y}, \bullet \sim N(\hat{\tau}_0, K_{\tau_0}^{-1}), \quad (19)$$

$$K_{\tau_0} = \frac{1}{b_{0,\tau}} + \frac{1}{e^{g_1}}, \quad \hat{\tau}_0 = K_{\tau_0}^{-1} \left(\frac{a_{0,\tau}}{b_{0,\tau}} + \frac{\tau_1}{e^{g_1}} \right).$$

Posteriors of h and g

To derive the posteriors of $\mathbf{h}$, we use the commonly employed approximate discrete mixture sampler of Kim et al. (1998). Since these manipulations are standard, we refer to Kim et al. (1998) for the conditional posterior of the states.

Posterior of λ and ν

$$\lambda_t \sim \mathcal{G}^{-1} \left(\frac{\nu + 1}{2}, \frac{1}{2} \frac{(y_t - \tau_t - \theta' Z_t)^2}{\exp(h_t)} \right) \quad (20)$$

$$\begin{aligned} p(\nu | \bullet) &\propto p(\boldsymbol{\lambda} p(\nu)) \\ &\propto \prod_{t=1}^T \frac{(\nu/2)^{\frac{\nu}{2}}}{\Gamma(\nu/2)} \lambda_t^{-(\frac{\nu}{2}+1)} e^{-\frac{\nu}{2\lambda_t}} \\ &= \frac{(\nu/2)^{\frac{T}{2}}}{\Gamma(\nu/2)^T} \left(\prod_{t=1}^T \right)^{-(\frac{\nu}{2}+1)} e^{\frac{\nu}{2} \sum_{t=1}^T \lambda_t^{-1}}. \end{aligned} \quad (21)$$

To sample from this distribution, we make use of an independent Metropolis-Hastings within Gibbs sampling step. The proposal distribution is defined as a normal with mean equal to the mode of 21 and covariance equal to the negative Hessian evaluated at the mode. To find the mode, we use the Newton-Raphson method. The Hessian is analytically available (Chan, 2017).

A.1.1 Further Insight into the GIGG Posterior

In order to highlight the effect of different choices on a_k and b_k , we write the posterior of the MIDAS coefficients in terms of its shrinkage coefficient representation:

$$\begin{aligned}\bar{\theta} &= \Lambda_*((\mathbf{Z}'\Lambda_{\lambda,h}^{-1}\mathbf{Z})^{-1} + \Lambda)\hat{\theta}_{k,j} \\ \bar{\theta}_{k,j} &\approx (1 - \kappa_{k,j})\hat{\theta}_{k,j}, \forall j = 1, \dots, L - r\end{aligned}\tag{22}$$

where $\tilde{\mathbf{y}} = \mathbf{y} - \boldsymbol{\tau}$ and $\hat{\theta}_{k,j} = (Z_{k,j}'\Lambda_{\lambda,h}^{-1}Z_{k,j})^{-1}Z_{k,j}'\Lambda_{\lambda,h}^{-1}\tilde{\mathbf{y}}$ can be viewed as a conditional maximum likelihood estimate for $\theta_{k,j}$. We suppress the indication of monthly frequency (m) for convenience here. Note that for this representation, we have assumed that the maximum likelihood estimate for group k exists and that the $L - r$ lags within $\mathbf{Z}_k$ have been orthogonalised. Since the group-size is likely to be much smaller than the sample size, and given that we will group-orthogonalise the data anyway (see section 2.4), these only represent very mild assumptions. Under these assumptions, it is easy to verify that the shrinkage coefficient $\kappa_{k,j|\bullet} = \frac{1}{1 + \tilde{s}_{k,j}^2 \vartheta^2 \gamma_k^2 \varphi_{k,j}^2}$ is bounded between 0 and 1 and thus dictates how far away the prior shrinks the coefficients from the maximum likelihood solution. It is easy to see that $\vartheta^2 \gamma_k^2 \varphi_{k,j}^2 \rightarrow \infty$, $\bar{\theta}_{k,j} \rightarrow \hat{\theta}_{k,j}$. The distribution $\pi(\kappa_{k,j})$ which is implicitly defined via the priors on γ_k^2 and $\varphi_{k,j}^2$ determine the a-priori shrinkage behaviour we can expect. By assuming $\gamma_k^2 \varphi_{k,j}^2 \sim \beta'(a_k, b_k)$, the joint distribution for $\boldsymbol{\kappa}_k$ can be factored as:

$$\begin{aligned}\pi(\kappa_k|\bullet) &= \frac{\Gamma(a_k + (L - r)b_k)}{\Gamma(a_k)\Gamma(b_k)^{k+1}} \prod_{j=1}^{p_k} \tilde{s}_{k,j}^{b_k} \left(1 + \sum_{j=1}^{p_k} \tilde{s}_{k,j} \frac{\kappa_{k,j}}{(1 - \kappa_{k,j})}\right)^{-(a_k + (1 + p_k)b_k)} \times \\ &\quad \left(\prod_{j=1}^{p_k} \kappa_{k,j}^{b_k - 1} (1 - \kappa_{k,j})^{-(b_k + 1)}\right),\end{aligned}\tag{23}$$

where $\tilde{s}_{k,j} = s_{k,j} \sum_{t=1}^T \frac{1}{\lambda_t} e^{-h_t}$ and $s_{k,j}$ is the j th lag's variance. This joint distribution factors into a dependent part, influenced by a_k , and an independent part, determined by b_k .

A.2 Group-Selection Algorithm

This section gives further details on the derivation of group-variable selection algorithm in 8. For convenience, we replicate the objective function here again:

$$\mathcal{L}(\tilde{\mathbf{Y}}, \alpha) = \frac{1}{2} \|\mathbf{Z}\alpha - \tilde{\mathbf{Y}}\|_2^2 + \sum_{k=1}^K \phi_k \|\alpha_k\|_2,\tag{24}$$

For simplicity, assume that the predictions $\tilde{\mathbf{Y}}$ can be decomposed as $\tilde{\mathbf{Y}} = \mathbf{Z}\theta + \tilde{\boldsymbol{\epsilon}}^y$, $\tilde{\boldsymbol{\epsilon}} \sim N(0, \Sigma)$. Intuitively, objective function 8 pushes those α_k to zero which have little influence on the predictions of our model, $\tilde{\mathbf{Y}}$.

We take the expectation with respect to 1) the expected risk and 2) $\theta|\mathbf{y}$ to account for all sources of uncertainty of the model^{B1}:

^{B1} Since α is independent of Σ , the integration of posterior uncertainty of Σ results in a constant and thus

$$\begin{aligned}
\mathcal{L}(\tilde{\mathbf{Y}}, \alpha) &= E_{\tilde{\mathbf{Y}}|\bullet}[\tilde{\mathbf{Y}}, \alpha|\bullet] \\
&= \sum_{k=1}^K \phi_k \|\alpha_k\|_2 + \frac{1}{2} \|\mathbf{Z}\alpha - \mathbf{Z}\theta\|_2^2 + \frac{1}{2} \text{tr}(\Sigma)
\end{aligned} \tag{25}$$

Then, taking the expectation with respect to $\theta|\mathbf{y}$:

$$\begin{aligned}
\mathcal{L}(\theta, \alpha) &= E_{\theta|\mathbf{y}}[\mathcal{L}(\tilde{\mathbf{Y}}, \alpha)] \\
&= \frac{1}{2} \|\mathbf{Z}\alpha - \mathbf{Z}\bar{\theta}\|_2^2 + \frac{1}{2} \text{tr}(\mathbf{Z}'\mathbf{Z}\Sigma_\theta) + \sum_{k=1}^K \phi_k \|\alpha_k\|_2,
\end{aligned} \tag{26}$$

where $\Sigma_\theta = \text{cov}(\theta)$ and $\bar{\theta} = E(\theta)$. Dropping all constant terms, the objective function reduces to:

$$\mathcal{L}(\theta, \alpha) = \frac{1}{2} \|\mathbf{Z}\alpha - \mathbf{Z}\bar{\theta}\|_2^2 + \sum_{k=1}^K \phi_k \|\alpha_k\|_2. \tag{27}$$

Notice, that we follow Chakraborty et al. (2020); Huber et al. (2021) by solving 26 on a Gibbs iteration bases (instead over the average of the posterior). Traditional solution methods such as the coordinate descent (Friedman et al., 2010) iteratively solve the sub-gradients L times for each group k until convergence:

$$\alpha_k^l = (\|r_k^{(l-1)}\|_2^2 - \phi_k)_+ \frac{r_k^{(l-1)}}{\|r_k^{(l-1)}\|_2^2}, \quad l = 1, \dots, L, \tag{28}$$

where r_k is the partial residual based on the previous iteration, $r^{(l)} = \mathbf{Z}'_k(\mathbf{y} - \mathbf{Z}_{-k}\alpha_{-k}^{(l-1)})$ and $-k$ refers to all but the k^{th} group. Orthonormalising $\mathbf{Z}_k$ via its SVD decomposition (which can conveniently be adapted as in (Breheny and Huang, 2015) when $L - r \gg T$), and stopping the coordinate descent after one iteration as per Ray and Bhattacharya (2018); Chakraborty et al. (2020) such that $\|r_k\| = \|\theta_k\|_2^2$ results immediately in 9.

B Additional results

B.1 Data

does not further influence the optimisation problem

Table B1: Macroeconomic indicators

Name	Acronym	Frequency	Transformation	Source
CBI: Volume of exp. Trades	CBI-ES	m	0	FAME
CBI: Volume of rep. Sales	CBI-S	m	0	FAME
CBI: Volume of exp. Output	CBI-EO	m	0	FAME
PMI: Manufacturing	PMI-M	m	0	FAME
PMI: Services	PMI-S	m	0	FAME
PMI: Construction	PMI-C	m	0	FAME
GfK Cons. Confidence	GfK	m	0	FAME
Index of Production	IoP	m	3	FAME
Index of Services	IoS	m	3	FAME
Exports	Exp	m	3	FAME
Imports	Imp	m	3	FAME
Unemployment Rate	UR	m	0	FAME
Employment	Emp	m	3	FAME
Job Vacancies	Vacancies	m	3	FAME
Hours Worked	Hours	m	3	FAME
Mortgage Approvals	Mortgage	m	3	FAME
VISA consumer spending	VISA	m	3	FAME
Real quarterly GDP growth (qoq)	GDP	q	3	FAME

Notes: The table shows the data used for the empirical application along with respective sampling frequencies (m = monthly, q = quarterly), tranformation applied (0 = no transformation, 1 = logs , 2 = first difference, 3 = growth rates) and data source. All data are downloaded from the UK data provider FAME. Please see the FAME website for further details on the data.

B.2 In-Sample Description

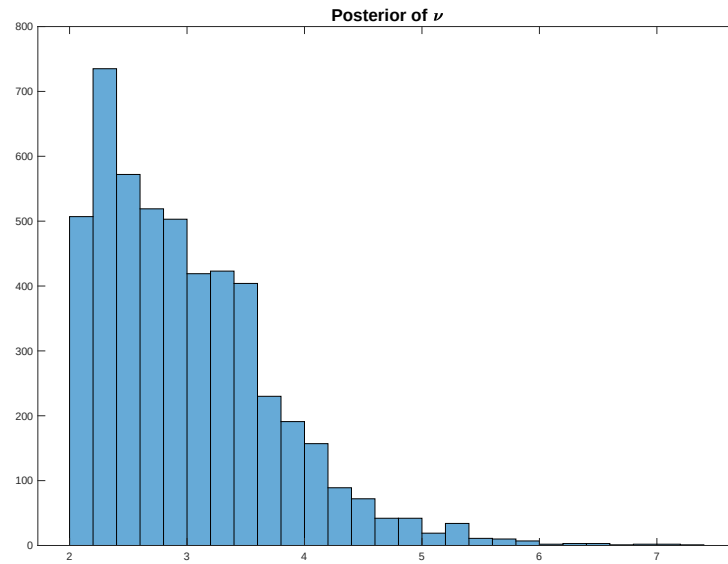


Figure B1: Posterior Degrees of Freedom of the t-distribution, based on the entire in-sample and last nowcast period.

Figure B1 has large mass over small posterior degrees of freedom for the t-distribution assumed for the errors of the observation equation of model [reference of equation 1]. The smaller the degrees of freedom, the more leptokurtic the tails of the observation equation's error distribution. Figure B1 shows strong identification of fat-tails.

B.3 Alternative hyperparameters on GIGG prior

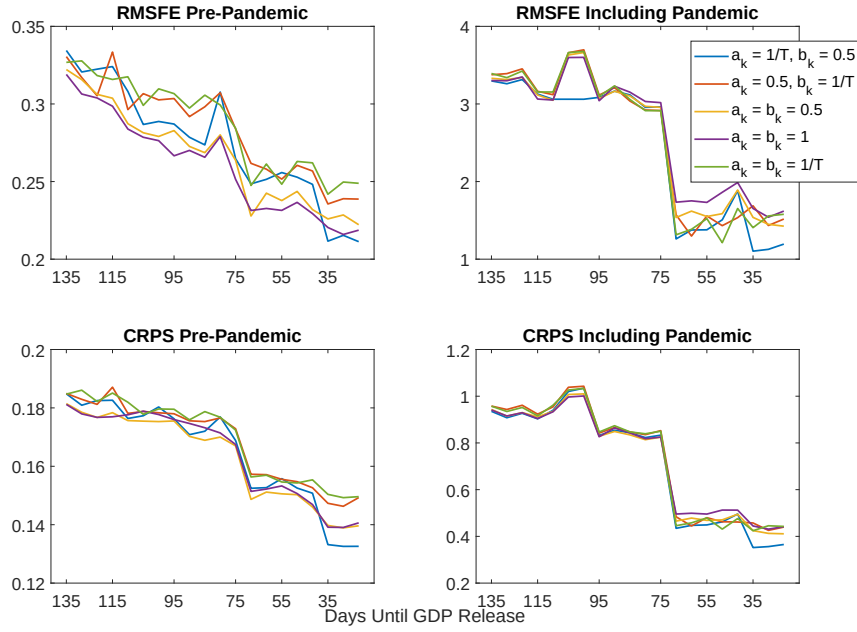


Figure B2: RSMFE and CRPS pre-pandemic and including pandemic for the Trend-SV-t-GIGG model with different hyper-parameter combinations. $1/T$ is adjusted to the in-sample length at each quarter the nowcasts are conducted. The GIGG prior implemented in the empirical application has $a_k = 1/T, b_k = 0.5$.

B.4 Trend-Cycle Decomposition and prior choice

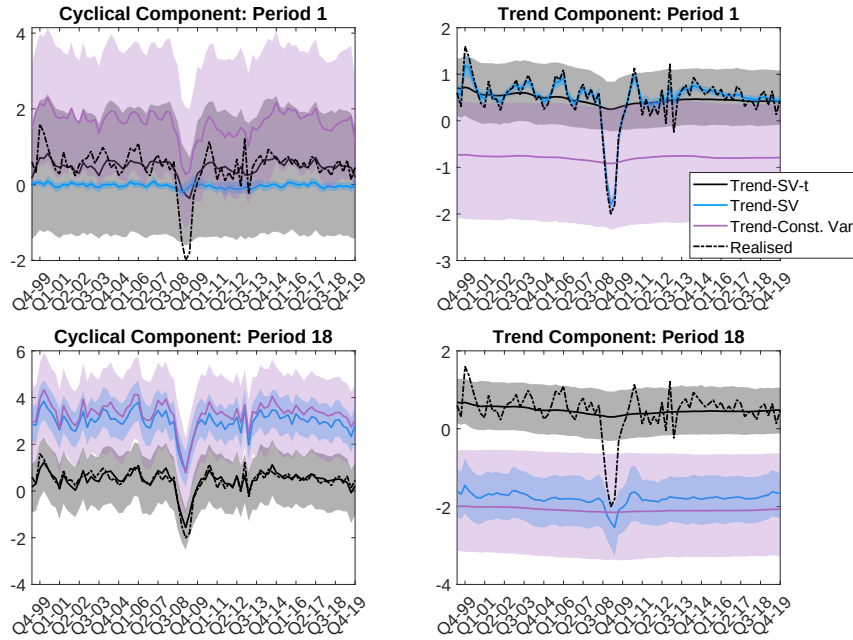


Figure B3: Posterior trend and cycle estimates, with **high** prior state variances.

Notes: State prior variances set to 1, prior on MIDAS coefficients as in baseline. Posterior means of trend up until 2019Q4 from the T-SV-t BMIDAS model (black), the T-SV (blue), and T-Const.Var (purple). Estimation at the 1st and 18th nowcast periods. Black dashed line: realisation of UK GDP growth. Shaded areas represent 95% credible intervals.

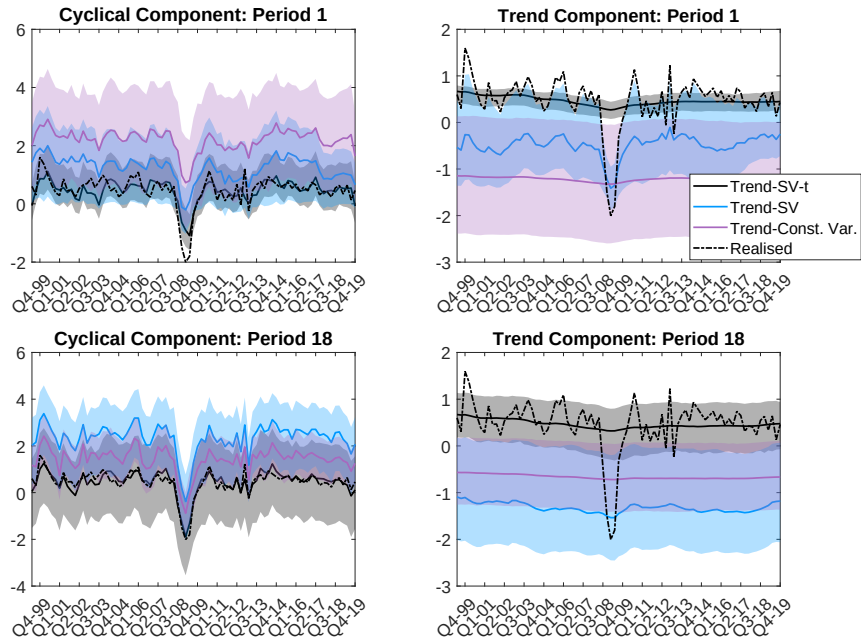


Figure B4: Posterior trend and cycle estimates, with **low** prior state variances.
Notes: State prior variances at 0.0001. MIDAS coeff. prior as in baseline. See notes Figure B3.

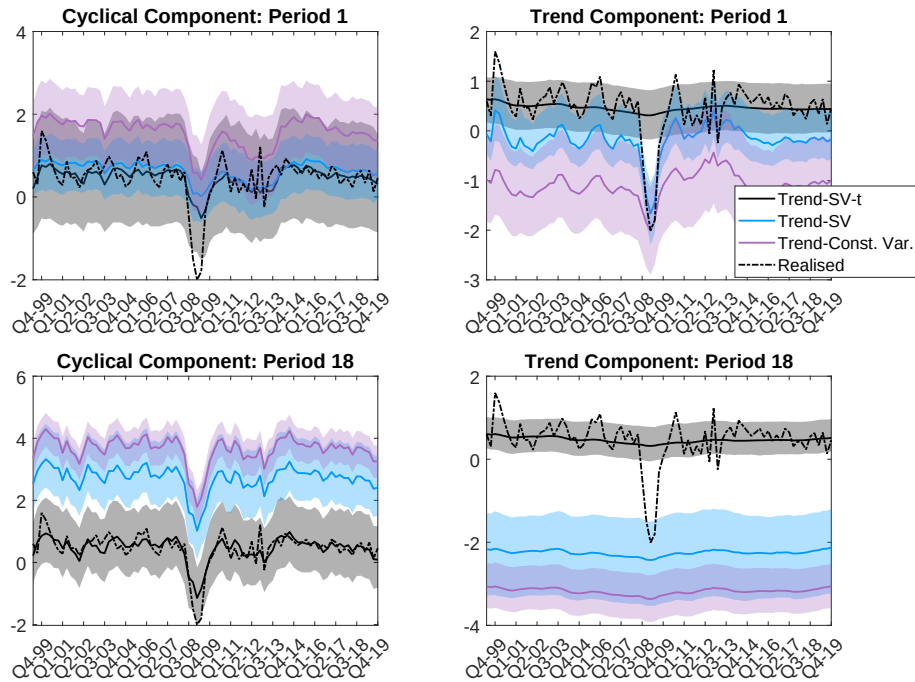


Figure B5: Posterior trend estimates using the horseshoe prior on MIDAS coefficients.
Notes: State prior variances as in baseline. Also see notes of Figure B3.

B.5 Alternatives to multivariate BMIDAS

Figure B6 illustrates the role of the BMIDAS component for nowcast performance against alternative ways of exploiting the information from high-frequency macroeconomic indicators, by comparing the T-SV-t BMIDAS with the following specifications (all models with a time-varying trend and stochastic volatility with t-distributed errors, with priors as in 2.2.2):

- *U-BMIDAS*: as baseline, but without Almon lag polynomial restrictions.
- *Combined univariate MIDAS*: MIDAS regressions for each of the K indicators, with a Normal prior, combined according to discounted RMSFE and CRPS (Stock and Watson, 2004), discount factor $\delta = 0.95$ (Carriero et al., 2019).
- *MF-DFM*: A mixed-frequency dynamic factor model that summarises co-movement across the K macroeconomic indicators contemporaneously and at up to two lags via a dynamic factor, akin to Antolin-Diaz et al. (2021).

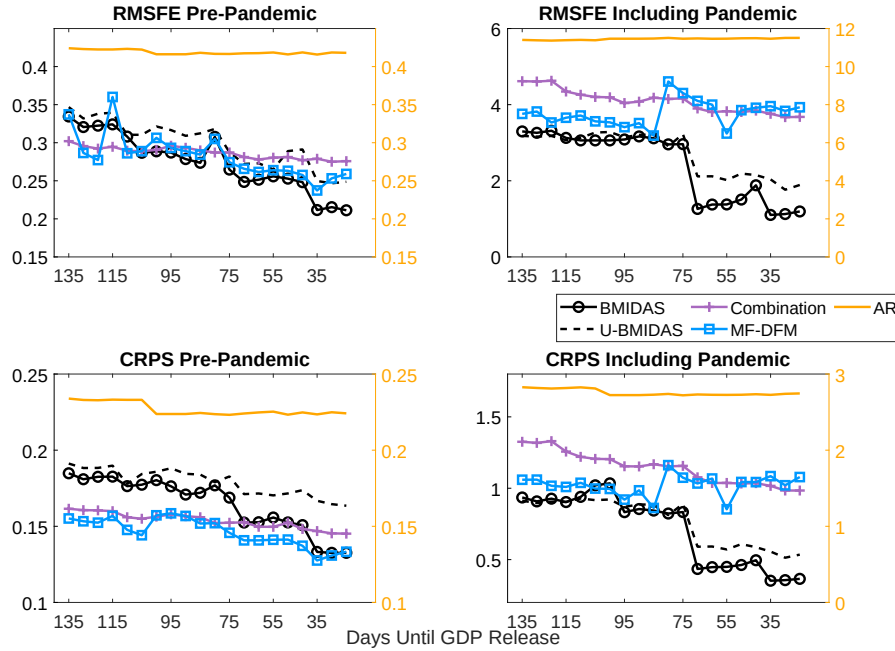


Figure B6: Nowcast performance against alternatives to BMIDAS specifications.

Notes: Absolute RMSFE and CRPS over nowcast periods (days before GDP release) for baseline BMIDAS (black solid, circle marker), U-BMIDAS (dashed black), Combined univariate MIDAS (purple, square), MF-DFM (blue, plus). Right y-axis: AR(2) (yellow). All with time-varying trend, SV, t-distributed errors.

Prior to the pandemic, point nowcast performance is quite similar across models. For density forecasts, the alternative DFM and combined univariate MIDAS outperform the BMIDAS specifications during the first half of the data release cycle somewhat. This is reflected in somewhat larger and more volatile uncertainty bands around the nowcasts of the BMIDAS model, as depicted in Figure B7. At the same time, the point and density nowcast performance of the BMIDAS model is continuously improving as new data get released. It matches or outperforms

the alternatives, but only once the “hard” economic data for the second month of the reference quarter are released, 35 days prior to GDP. The U-BMIDAS does not show the same extent of improvement, such that regularising the coefficients via Almon lag coefficients, in addition to the Bayesian shrinkage, appears to help exploiting incoming information. When including the pandemic into the evaluation period, both BMIDAS models clearly outperform the DFM and the combined univariate MIDAS. This reflects that these models capture the recovery after the pandemic early on, whereas alternative models almost entirely miss it. At nowcast period 13, 65 days prior to GDP release when “hard” economic data for the first month of the reference quarter are released, both MIDAS models, but in particular the proposed Almon polynomial restricted model strongly improve in point and density nowcast performance. At this stage, these models also update their nowcast on the trough, capturing it almost fully. The DFM’s performance even slightly worsens at this stage, while the combined univariate MIDAS improves only gradually and slightly as new information comes in.

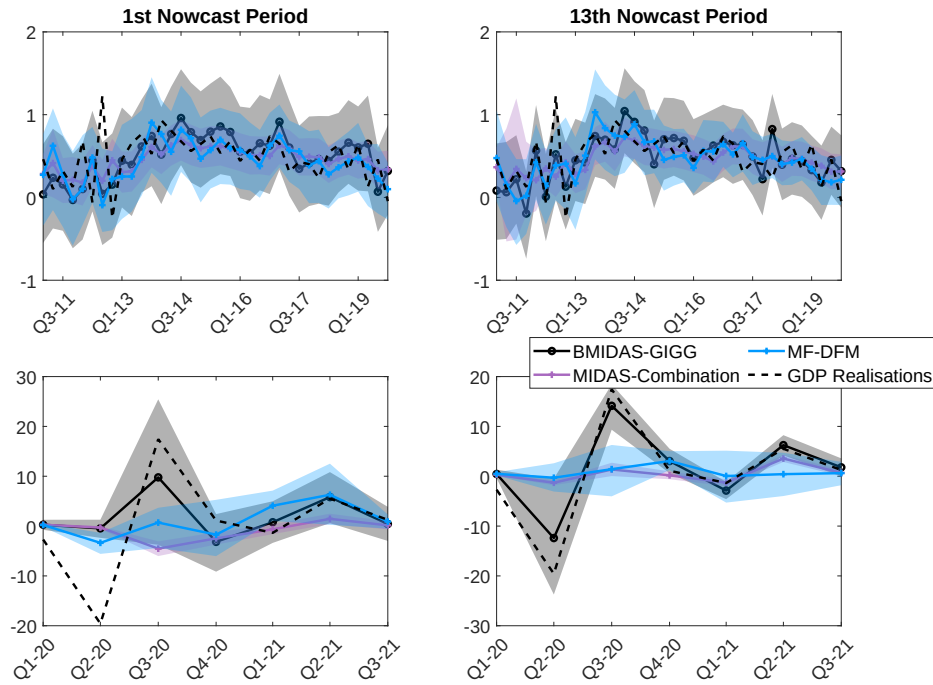


Figure B7: Posterior mean and density nowcasts, for alternative MIDAS specifications.

Notes: Posterior mean of nowcasts for 2011Q1 to 2019Q4, 1st and 13th nowcast periods, and over the pandemic quarters (lower panel). Models are the same as in Figure B6. Shaded areas show 95% credible intervals. Black dashed lines show quarterly GDP growth realisations.

B.6 U-BMIDAS vs BMIDAS with Almon lag restrictions

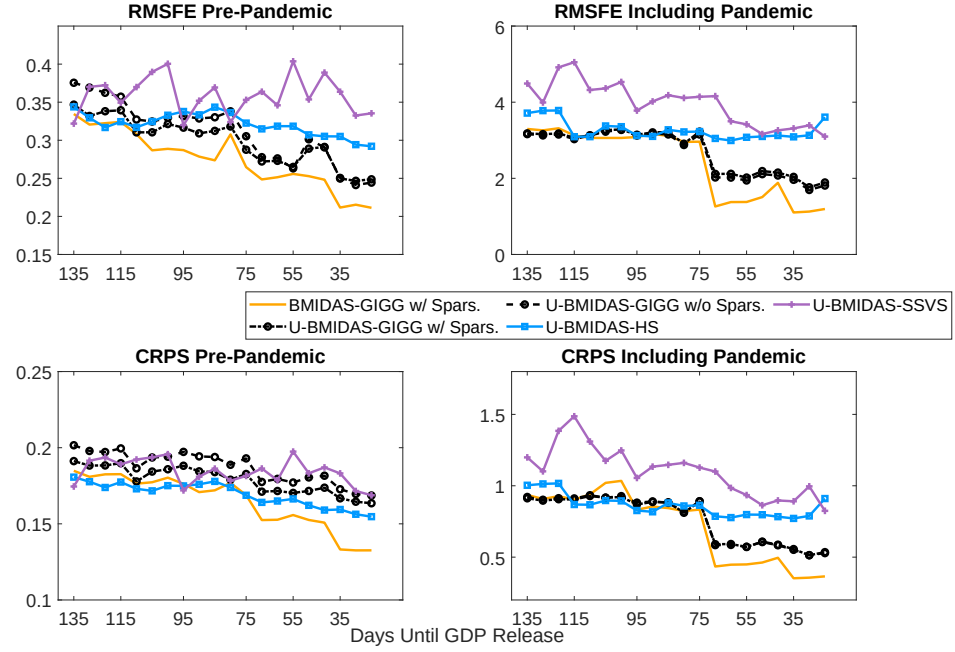


Figure B8: Nowcast performance against alternative priors on U-MIDAS regressions.

Notes: RMSFE and CRPS for baseline BMIDAS (yellow), and U-MIDAS with alternative priors. All models with time-varying trend and stochastic volatility with t-distributed errors.

B.7 Inclusion Probabilities

Table B2: Posterior Inclusion Probabilities

Nowcast Periods	Evaluation pre-pandemic				Evaluation incl. pandemic period			
	Average	6	13	18	Average	6	13	18
GIGG prior								
PMI-M	0.32 (0.20)	0.28 (0.20)	0.21 (0.17)	0.19 (0.15)	0.32 (0.19)	0.46 (0.24)	0.15 (0.12)	0.13 (0.12)
PMI-S	0.25 (0.19)	0.30 (0.21)	0.22 (0.16)	0.19 (0.15)	0.53 (0.18)	0.65 (0.20)	0.19 (0.16)	0.61 (0.24)
IoP	0.10 (0.09)	0.06 (0.06)	0.10 (0.09)	0.27 (0.19)	0.36 (0.21)	0.32 (0.20)	0.29 (0.20)	0.62 (0.22)
IoS	0.44 (0.13)	0.24 (0.18)	0.81 (0.15)	0.93 (0.06)	0.73 (0.12)	0.48 (0.25)	0.97 (0.03)	0.99 (0.01)
UE	0.28 (0.19)	0.31 (0.21)	0.24 (0.19)	0.12 (0.11)	0.29 (0.20)	0.36 (0.24)	0.18 (0.15)	0.14 (0.12)
Vacancies	0.14 (0.11)	0.12 (0.10)	0.13 (0.11)	0.12 (0.10)	0.30 (0.20)	0.29 (0.20)	0.18 (0.15)	0.27 (0.19)
Mortgage	0.09 (0.08)	0.10 (0.09)	0.08 (0.07)	0.05 (0.05)	0.42 (0.19)	0.51 (0.19)	0.35 (0.19)	0.28 (0.18)
HS prior								
PMI-M	0.49 (0.21)	0.41 (0.24)	0.55 (0.24)	0.72 (0.20)	0.64 (0.19)	0.74 (0.16)	0.48 (0.25)	0.75 (0.16)
PMI-S	0.22 (0.15)	0.16 (0.15)	0.02 (0.02)	0.33 (0.22)	0.45 (0.15)	0.64 (0.17)	0.02 (0.02)	0.54 (0.22)
IoP	0.09 (0.07)	0.21 (0.17)	0.02 (0.02)	0.05 (0.05)	0.32 (0.17)	0.53 (0.23)	0.07 (0.06)	0.29 (0.20)
IoS	0.10 (0.07)	0.01 (0.01)	0.26 (0.18)	0.04 (0.04)	0.47 (0.09)	0.05 (0.05)	0.83 (0.04)	0.80 (0.05)
UE	0.03 (0.03)	0.00 (0.00)	0.02 (0.02)	0.03 (0.04)	0.03 (0.03)	0.00 (0.00)	0.02 (0.01)	0.06 (0.05)
Vacancies	0.15 (0.11)	0.04 (0.04)	0.29 (0.20)	0.27 (0.19)	0.30 (0.13)	0.15 (0.13)	0.29 (0.20)	0.60 (0.19)
Mortgage	0.21 (0.14)	0.11 (0.10)	0.32 (0.22)	0.32 (0.22)	0.51 (0.15)	0.54 (0.21)	0.60 (0.20)	0.77 (0.13)
SSVS prior								
PMI-M	0.32 (0.08)	0.04 (0.03)	0.32 (0.21)	0.91 (0.08)	0.28 (0.07)	0.05 (0.02)	0.19 (0.15)	0.87 (0.11)
PMI-S	0.14 (0.14)	0.03 (0.07)	0.00 (0.02)	0.46 (0.00)	0.23 (0.25)	0.48 (0.09)	0.00 (0.07)	0.49 (0.00)
IoP	0.01 (0.01)	0.00 (0.01)	0.00 (0.00)	0.01 (0.02)	0.02 (0.02)	0.03 (0.02)	0.00 (0.00)	0.05 (0.04)
IoS	0.01 (0.01)	0.01 (0.01)	0.02 (0.03)	0.00 (0.00)	0.42 (0.06)	0.02 (0.08)	0.82 (0.01)	0.70 (0.08)
UE	0.07 (0.02)	0.04 (0.05)	0.00 (0.00)	0.00 (0.00)	0.02 (0.01)	0.05 (0.06)	0.00 (0.01)	0.00 (0.00)
Vacancies	0.08 (0.06)	0.01 (0.02)	0.02 (0.02)	0.31 (0.21)	0.11 (0.07)	0.03 (0.03)	0.01 (0.01)	0.41 (0.23)
Mortgages	0.18 (0.09)	0.00 (0.00)	0.14 (0.12)	0.60 (0.24)	0.35 (0.11)	0.28 (0.16)	0.34 (0.20)	0.86 (0.12)

Notes: Mean posterior inclusion probability. Standard deviation estimates are in parentheses. All models with time-varying trend, stochastic volatility, and t-distributed errors.