

Online Appendix for

“On the Effects of Monetary Policy Shocks on Income and Consumption Heterogeneity”

Minsu Chang*

Seoul National University

Frank Schorfheide

University of Pennsylvania,

CEPR, PIER, NBER

July 8, 2026

This Appendix consists of the following sections:

- A. Density Estimation
- B. Prior and Posterior Computations
- C. State-Space Representation for Mixed-Frequency Analysis
- D. Data Used in the Empirical Analysis
- E. Knot Placement, Lag Length, and Hyperparameter Selection
- F. Additional Empirical Results

* Correspondence: M. Chang: Department of Economics, Seoul National University, Seoul, South Korea. Email: minsuchang@snu.ac.kr (Chang). F. Schorfheide: Department of Economics, University of Pennsylvania, Philadelphia, PA 19104-6297. Email: schorf@ssc.upenn.edu (Schorfheide). We thank Xiaohong Chen, Stephanie Ettmeier, Eva Janssens, Chi Hyun Kim, and participants at various seminars and conferences for helpful suggestions. This work was supported by the New Faculty Startup Fund from Seoul National University.

A Density Estimation

A.1 Top Coding

Likelihood Function with Censoring. We define the censoring point c_t as

$$c_t = \max_{i=1,\dots,N} x_{it}$$

Moreover, we let

$$N_{t,max} = \sum_{i=1}^N \mathbb{I}\{x_{it} = c_t\}.$$

If $N_{t,max} = 1$, we assume that the observed sample is not constrained by the top-coding and use the standard likelihood function described in the main text. If $N_{t,max} > 1$ we use a likelihood function that assumes that any earnings value exceeding c_t is coded as c_t .

Recall that in the main text we ignored the dependence of the cross-sectional sample size N on t in the notation and defined $p^{(K)}(X_t|\alpha_t) = \exp\{N\mathcal{L}^{(K)}(\alpha_t|X_t)\}$, where

$$\mathcal{L}^{(K)}(\alpha_t|X_t) = \bar{\zeta}'(X_t)\alpha_t - \ln \int_0^\infty \exp\{\zeta'(x)\alpha_t\}dx, \quad \bar{\zeta}(X_t) = \frac{1}{N} \sum_{i=1}^{N_t} \zeta(x_{it}).$$

We introduce the unknown parameter $\pi_t = \mathbb{P}\{x_{it} \geq c_t\}$. We drop the top-coded observations from the definition of $\bar{\zeta}(X_t)$ and make the time dependence explicit in the notation. Let

$$\bar{\zeta}_t(X_t) = \frac{1}{N_t} \sum_{i=1}^{N_t} \zeta(x_{it}) \mathbb{I}\{x_{it} < c_t\}. \quad (\text{A.1})$$

The log-likelihood function is obtained as follows: the sample contains $N_{t,max}$ top-coded observations where the probability of sampling a top-coded observation is π_t . The probability of sampling an observation that is not top-coded is $(1 - \pi_t)$. Conditional on not being top-coded, the observation $x_{it} < c_t$ is sampled from a continuous density with a domain that is truncated at c_t . Thus, dividing the log-likelihood by the sample size N_t , we obtain

$$\begin{aligned} \mathcal{L}^{(K)}(\alpha_t, \pi_t|X_t) &= \frac{N_{t,max}}{N_t} \ln \pi + \frac{N_t - N_{t,max}}{N_t} \ln(1 - \pi_t) \\ &\quad + \bar{\zeta}'_t(X_t)\alpha_t - \frac{N_t - N_{t,max}}{N_t} \ln \int_0^{c_t} \exp\{\zeta'(x)\alpha_t\}dx. \end{aligned} \quad (\text{A.2})$$

Notice that regardless of the value of α_t , the MLE of π_t is

$$\hat{\pi}_t = \operatorname{argmax}_{\pi \in [0,1]} \mathcal{L}^{(K)}(\alpha_t, \pi_t|X_t) = N_{t,max}/N_t. \quad (\text{A.3})$$

Moreover, regardless of the value of π_t , the MLE of α_t is given by

$$\begin{aligned}\hat{\alpha}_t &= \operatorname{argmax}_{\alpha_t} \mathcal{L}^{(K)}(\alpha_t, \pi_t | X_t) \\ &= \operatorname{argmax}_{\alpha_t} \bar{\zeta}'_t(X_t) \alpha_t - \frac{N_t - N_{t,max}}{N_t} \ln \int_0^{c_t} \exp \{ \zeta'(x) \alpha_t \} dx.\end{aligned}\tag{A.4}$$

The objective function for α_t is almost identical to what we had without top coding, except for a definition of $\bar{\zeta}_t(X_t)$ that drops the top-coded observations in the summation and the factor of $(N_t - N_{t,max})/N_t$ in front of the normalization constant of the density.

Recovering the Density for Uncensored Observations. To reconstruct the full density we can use

$$p(x|\alpha_t) = \frac{\exp \left\{ \sum_{k=1}^K \alpha_{k,t} \zeta_k(x) \right\}}{\int_0^\infty \exp \left\{ \sum_{k=1}^K \alpha_{k,t} \zeta_k(x) \right\} dx}.\tag{A.5}$$

Note that here we drop the censoring indicator function and the integration is now from 0 to ∞ . Once the α_t 's have been estimated based on the censored observations, we work with the full density in the functional state-space model and its K -dimensional approximation.

Modification of Hessian Matrix. We now re-compute the score and the Hessian. Dropping the (K) superscript we obtain the following first derivatives with respect to α_k for $k = 1, \dots, K$:

$$\mathcal{L}_k^{(1)}(\alpha_t | \pi_t, X_t) = \bar{\zeta}_{t,k}(X_t) - \left(\frac{N_t - N_{t,max}}{N_t} \right) \int_0^{c_t} \zeta_k(x) \bar{p}(x|\alpha_t) dx,$$

where

$$\bar{p}(x|\alpha_t) = \frac{\exp \left\{ \sum_{k=1}^K \alpha_{k,t} \zeta_k(x) \right\}}{\int_0^{c_t} \exp \left\{ \sum_{k=1}^K \alpha_{k,t} \zeta_k(x) \right\} dx} \mathbb{I}\{x < c_t\}.$$

We can now deduce from our previous calculations that

$$\begin{aligned}\mathcal{L}_{kl}^{(2)}(\alpha_t | \pi_t, X_t) &= - \left(\frac{N_t - N_{t,max}}{N_t} \right) \int_0^{c_t} \left(\zeta_k(x) - \int_0^{c_t} \zeta_k(x) \bar{p}(x|\alpha_t) dx \right) \\ &\quad \times \left(\zeta_l(x) - \int_0^{c_t} \zeta_l(x) \bar{p}(x|\alpha_t) dx \right) \bar{p}(x|\alpha_t) dx.\end{aligned}\tag{A.6}$$

Thus, compared to the standard case, the limits of integration change and there is an additional factor $(N_t - N_{t,max})/N_t$.

A.2 Transformations of the $\hat{\alpha}_t$ s

Compression/Standardization. The vector $\hat{\hat{\alpha}}_t = \hat{\alpha}_t - \alpha_*$ may exhibit collinearity. Even though K basis functions may be necessary to approximate the cross-sectional densities, the time variation might be concentrated in a lower-dimensional space, because, for instance, only the means of the cross-sectional distributions are varying over time. This feature can be captured by assuming that the time-variation is captured by a $\tilde{K} < K$ dimensional factor a_t :

$$(\alpha_t - \alpha_*)' = a_t' \Lambda, \quad (\text{A.7})$$

where Λ is a $\tilde{K} \times K$ matrix. As is well known from the factor model literature, Λ and a_t are only identified up to a $\tilde{K} \times \tilde{K}$ dimensional invertible matrix. In principle, the matrix Λ and the sequence of vectors a_t , $t = 1, \dots, T$ have to be estimated simultaneously under this factor structure,

To avoid the simultaneous estimation of the cross-sectional densities, we take the following short cut. First, we compute the $\hat{\alpha}_t$ s period-by-period without imposing any restrictions. Second, conditional on α_* we compute the demeaned (and potentially seasonally adjusted) MLEs $\hat{\hat{\alpha}}_t = \hat{\alpha}_t - \alpha_*$ and arrange them in a $T \times K$ matrix $\hat{\hat{\alpha}}$ with rows $\hat{\hat{\alpha}}_t'$. Third, we conduct a principal components analysis which is based on the eigenvalue decomposition of the sample covariance matrix $\hat{\hat{\alpha}}' \hat{\hat{\alpha}} / T$. Let $\hat{M}$ be $K \times \tilde{K}$ matrix of eigenvectors associated with the $\tilde{K}$ non-zero eigenvalues (in practice greater than 10^{-10}).¹ Then, let

$$\hat{a} = \hat{\hat{\alpha}} \hat{M}, \quad \hat{\Lambda} = (\hat{a}' \hat{a})^{-1} \hat{a}' \hat{\hat{\alpha}}, \quad (\text{A.8})$$

where $\hat{a}$ is the $T \times \tilde{K}$ matrix with rows $\hat{a}_t'$. Even if $\tilde{K} = K$ this operation standardizes the basis function coefficients α_t . To evaluate the MDD formula (28), we replace K by $\tilde{K}$, $\mathcal{L}^{(K)}(\hat{\alpha}_t | X_t)$ by $\mathcal{L}^{(\tilde{K})}(\alpha_* + \hat{\Lambda}' \hat{a}_t | X_t)$, and we change the term $\sum_{t=1}^T \ln |\hat{V}_t|^{1/2}$ to $\sum_{t=1}^T \ln |(\hat{\Lambda} \hat{V}_t^{-1} \hat{\Lambda}')^{-1}|^{1/2}$.

Seasonal Adjustments. Deterministic seasonal adjustments of the cross-sectional densities can be incorporated in the model, as needed, by replacing the vector of constants $\alpha_* = \alpha_t - \tilde{\alpha}_t$ by a time-varying process. In our application the time period t is either a quarter or a month. For quarterly data we let $\alpha_{*,t} = \sum_{q=1}^4 \alpha_{q,t} s_q(t)$, where $s_q(t) = 1$ if period t is associated with quarter q and $s_q(t) = 0$ otherwise. For monthly data we use $\alpha_{*,t} = \sum_{m=1}^{12} \alpha_{m,t} s_m(t)$, where $s_m(t) = 1$ if period t is associated with month m and $s_m(t) = 0$ otherwise.

¹Because our goal is to eliminate perfect collinearities, we choose an eigenvalue cut-off that yields $\alpha_* + \hat{\Lambda}' \hat{a}_t = \hat{\hat{\alpha}}_t$.

A.3 Recovering Cross-Sectional Densities

Based on the estimated state-transition equation we can generate forecasts and impulse response functions for the compressed coefficients a_t . However, the dynamics of these coefficients in itself are not particularly interesting. Thus, we have to convert them back into densities using the following steps (which can be executed for each prior/posterior draw of a_t from the relevant posterior distribution). First, use (A.7) with $\Lambda = \hat{\Lambda}$ to transform a_t into α_t . If the estimation is based on a seasonal adjustment, α_* can be replaced by $\alpha_{*,t}$, or, if the goal is to compute impulse responses, one could use the average of the seasonal dummies as intercept. Second, compute

$$p^{(K)}(x|\alpha_t) = \frac{\exp\{\zeta'(x)\alpha_t\}}{\int \exp\{\zeta'(\tilde{x})\alpha_t\}d\tilde{x}}.$$

B Prior and Posterior Computations

B.1 More Details on the Prior

Recall that the i th equation of the W_t VAR is given by (22). We assume that the parameters (β_i, D_i) are *a priori* independent across equations, i.e.,

$$p(\beta, D) = \prod_{i=1}^n p(\beta_i | D_i) p(D_i). \quad (\text{A.9})$$

For each pair (β_i, D_i) we use a Normal-Inverse Gamma (NIG) distribution of the form

$$\beta_i | D_i \sim \mathcal{N}(\underline{\beta}_i, D_i \underline{V}_i^\beta), \quad D_i \sim IG(\underline{\nu}_i, \underline{S}_i). \quad (\text{A.10})$$

The prior density takes the form

$$\begin{aligned} p(\beta_i, D_i) &= (2\pi)^{-k_i/2} |\underline{V}_i^\beta|^{-1/2} \frac{\underline{S}_i^{\underline{\nu}_i}}{\Gamma(\underline{\nu}_i)} D_i^{-(\underline{\nu}_i+1+k_i/2)} \\ &\quad \times \exp \left\{ -\frac{1}{D_i} \left[\underline{S}_i + \frac{1}{2} (\beta_i - \underline{\beta}_i)' (\underline{V}_i^\beta)^{-1} (\beta_i - \underline{\beta}_i) \right] \right\}. \end{aligned}$$

In the remainder of this subsection we discuss the construction of $\underline{\beta}_i$, $\underline{V}_i^\beta$, $\underline{\nu}_i$, and $\underline{S}_i$. The prior is obtained by transforming a prior for the reduced-form parameters (Φ, Σ) into a prior for the quasi-structural parameters $(\beta_1, \dots, \beta_{n_w}, D)$.

Prior for D_i and the α_i component of β_i . We start from a prior for $\Sigma = A^{-1'}DA^{-1}$:

$$\Sigma \sim IW(\underline{\nu}, \underline{S}), \quad \underline{S} = \text{diag}(\underline{s}_1^2, \dots, \underline{s}_n^2). \quad (\text{A.11})$$

Chan (2021) shows that this prior implies

$$D_i \sim IG\left(\frac{\underline{\nu} + i - n}{2}, \frac{\underline{s}_i^2}{2}\right), \quad i = 1, \dots, n. \quad (\text{A.12})$$

Thus, a comparison with (A.10) indicates that we are setting

$$\underline{\nu}_i = \frac{\underline{\nu} + i - n}{2}, \quad \underline{S}_i = \frac{\underline{s}_i^2}{2}. \quad (\text{A.13})$$

Moreover, (A.11) implies that

$$A_{ij}|D_i \sim \mathcal{N}\left(0, \frac{D_i}{\underline{s}_j^2}\right), \quad 1 \leq j < i, \quad i = 2, \dots, n, \quad (\text{A.14})$$

which determines the prior for the α_i component of β_i .

Prior for the $B_{.i}$ component of β_i . Recall that $B_{.i}$ consists of np coefficients on lagged elements of y_t and an intercept. The overall dimension of the vector is $k \times 1$. The prior will take the form

$$B_{.i} \sim \mathcal{N}(\underline{B}_{.i}, \underline{V}_{.i}^B), \quad (\text{A.15})$$

where the $k \times k$ matrix $\underline{V}_{.i}^B$ is assumed to be diagonal.

To specify a prior for $B_{.i}$, we loosely map *a priori* beliefs about $(\alpha_i, \Phi_{.i})$ into beliefs about $B_{.i}$. To simplify the notation a bit, let $\phi_i = \Phi_{.i}$ and suppose that the researcher starts with the belief that

$$\phi_i \sim \mathcal{N}(\underline{\phi}_i, D_i \underline{V}_i^\phi) \quad (\text{A.16})$$

with $\underline{\phi}_i = 0$. Because the macroeconomic variables are in log-level we let for $i = 1, \dots, n_y$:

$$[\underline{\phi}_i]_j = \begin{cases} 1 & \text{if } j = i \\ 0 & \text{otherwise} \end{cases} \quad j = 1, \dots, k. \quad (\text{A.17})$$

The $k \times k$ prior covariance matrix $\underline{V}_i^\phi$ is assumed to be diagonal with elements $l = 1, \dots, k$:

$$[\underline{V}_i^\phi]_{ll} = \begin{cases} \frac{1}{\lambda_1} \frac{1}{\underline{s}_i^2 h^{\lambda_4}} & \text{for coeff. on the } h\text{-th lag if vars } (i, j) \text{ belong to same block} \\ \frac{1}{\lambda_1} \frac{1}{\lambda_2 \underline{s}_i^2 h^{\lambda_4}} & \text{for coeff. on the } h\text{-th lag if var } i \text{ belongs to } Y \text{ and } j \text{ belongs to } a \\ \frac{1}{\lambda_1} \frac{1}{\lambda_3 \underline{s}_i^2 h^{\lambda_4}} & \text{for coeff. on the } h\text{-th lag if var } i \text{ belongs to } a \text{ and } j \text{ belongs to } Y \\ \frac{1}{\lambda_5} & \text{for the intercept} \end{cases}$$

We now turn the prior for ϕ_i into a prior for $B_{\cdot i}$, utilizing the relationship between the quasi-structural-form coefficients, B , and the reduced form coefficients, Φ . For the coefficients on lagged elements of y_t we obtain:

$$[B_h]_{ij} = [\Phi_h]_{ij} + \sum_{l=1}^{i-1} A_{il}[\Phi_h]_{lj}. \quad (\text{A.18})$$

Likewise, the $n \times 1$ vector of intercepts, B_0 , is related to the reduced form intercept via

$$[B_0]_i = [\Phi_0]_i + \sum_{l=1}^{i-1} A_{il}[\Phi_0]_l. \quad (\text{A.19})$$

Taking expectations of (A.18) and (A.19), and using $\mathbb{E}[A_{ij}] = 0$, we deduce that

$$\mathbb{E}[[B_h]_{ij}] = \mathbb{E}[[\Phi_h]_{ij}], \quad \mathbb{E}[[B_0]_i] = \mathbb{E}[[\Phi_0]_i] \quad (\text{A.20})$$

We use a prior covariance matrix $\underline{V}_i^B$ that is diagonal. The entries on the diagonal are specified as follows: we first express the variance of a generic element $[B_h]_{ij}$ in terms of the variances of A_{ij} and $[\Phi_h]_{ij}$:

$$\begin{aligned} \mathbb{V}[[B_h]_{ij}] &= \mathbb{E}[\mathbb{V}([B_h]_{ij})|A] + \mathbb{V}[\mathbb{E}([B_h]_{ij})|A] \\ &= \mathbb{E} \left[\mathbb{V}([\Phi_h]_{ij}) + \sum_{l=1}^{i-1} A_{il}^2 \mathbb{V}([\Phi_h]_{lj}) \right] + \mathbb{V} \left[\mathbb{E}([\Phi_h]_{ij}) + \sum_{l=1}^{i-1} A_{il} \mathbb{E}([\Phi_h]_{lj}) \right] \\ &= \mathbb{V}([\Phi_h]_{ij}) + \sum_{l=1}^{i-1} \mathbb{V}(A_{il}) \mathbb{V}([\Phi_h]_{lj}) + \sum_{l=1}^{i-1} \mathbb{V}(A_{il}) (\mathbb{E}([\Phi_h]_{lj}))^2. \end{aligned}$$

The same calculation for the variance of the intercept leads to:

$$\begin{aligned} \mathbb{V}[[B_0]_i] &= \mathbb{E}[\mathbb{V}([B_0]_i)|A] + \mathbb{V}[\mathbb{E}([B_0]_i)|A] \\ &= \mathbb{E} \left[\mathbb{V}([\Phi_0]_i) + \sum_{l=1}^{i-1} A_{il}^2 \mathbb{V}([\Phi_0]_l) \right] + \mathbb{V} \left[\mathbb{E}([\Phi_0]_i) + \sum_{l=1}^{i-1} A_{il} \mathbb{E}([\Phi_0]_l) \right] \\ &= \mathbb{V}([\Phi_0]_i) + \sum_{l=1}^{i-1} \mathbb{V}(A_{il}) \mathbb{V}([\Phi_0]_l) + \sum_{l=1}^{i-1} \mathbb{V}(A_{il}) (\mathbb{E}([\Phi_0]_l))^2. \end{aligned}$$

To arrange the first np $\mathbb{V}[[B_h]_{ij}]$ terms on the diagonal of the $k \times k$ matrix $\underline{V}_i^B$ we use the index function

$$f(j, h) = (h - 1)n + j. \quad (\text{A.21})$$

Here h corresponds to the lag and j is the index for the element of the y_{t-h} vector. Using the definition of the index function, the expressions for $\mathbb{V}[A_{il}]$ and $\mathbb{E}([\Phi_h]_{lj})$ from the Normal

distribution in (A.14), and the expression for $\mathbb{V}([\Phi_h]_{lj})$ from the Normal distribution in (A.16), we can write

$$\mathbb{V}([B_h]_{ij}) = D_i[\underline{V}_i^\phi]_{f(j,h)} + \sum_{l=1}^{i-1} \frac{D_i}{\underline{s}_l^2} \left[D_l[\underline{V}_l^\phi]_{f(j,h)} + [\underline{\phi}_l]_{f(j,h)}^2 \right].$$

For the intercept in equation i we obtain

$$\mathbb{V}([B_0]_i) = D_i \frac{D_i}{\lambda_5} + \sum_{l=1}^{i-1} \frac{D_i D_l}{\underline{s}_l^2 \lambda_5}.$$

We now replace the variance parameter D_l by the hyperparameter $\underline{s}_l^2$. This ensures that $\underline{V}_i^B$ is not a function of the (unknown) variance parameter D_i and simplifies posterior calculations. Using

$$\mathbb{V}([B_h]_{ij}) = D_i[\underline{V}_i^B]_{f(j,h)} \quad \text{and} \quad \mathbb{V}([B_0]_i) = D_i[\underline{V}_i^B]_{np+1}$$

we obtain

$$[\underline{V}_i^B]_{f(j,h)} = [\underline{V}_i^\phi]_{f(j,h)} + \sum_{l=1}^{i-1} \left([\underline{V}_l^\phi]_{f(j,h)} + \frac{1}{\underline{s}_l^2} [\underline{\phi}_l]_{f(j,h)}^2 \right). \quad (\text{A.22})$$

$$[\underline{V}_i^B]_{np+1} = \frac{1}{\lambda_5} + \sum_{l=1}^{i-1} \frac{1}{\lambda_5} = \frac{i}{\lambda_5}. \quad (\text{A.23})$$

Summary. The overall prior takes the form (A.10). The prior for D_i is given by (A.12). The prior for β_i is obtained by combining (A.14) with (A.15), where mean and variance are given in (A.20) and (A.22), respectively. The hyperparameters for the prior are summarized in Table A-1. We set $\underline{s}_i$ equal to the sample standard deviation of W_i .

B.2 Posterior Sampling and MDD

Model and prior are set up so that the coefficients can be estimated equation by equation:

$$p(W, \beta, D) = \prod_{i=1}^N \left((2\pi D_i)^{-1/2} \exp \left\{ -\frac{1}{2D_i} (W_i - Z_i \beta_i)' (W_i - Z_i \beta_i) \right\} p(\beta_i | D_i) p(D_i) \right) \quad (\text{A.24})$$

Because the prior is conjugate, the posterior stays in the NIG family. It takes the form

$$\beta_i | (D_i, W_i) \sim \mathcal{N}(\bar{\beta}_i, D_i \bar{V}_i^\beta), \quad D_i \sim IG(\bar{\nu}_i, \bar{S}_i), \quad (\text{A.25})$$

Table A-1: Hyperparameters for VAR Prior

Parameter	Description
$\underline{\nu} = 2n$	Degrees of freedom for IG distribution
$\underline{s}_i = \text{StDev}(W_i)$	Shape para for IG; use sample standard dev.
λ_1	Overall precision of prior
λ_2	Relative precision for a to Y transmission
$\lambda_3 = 1$	Relative precision for Y to a transmission
$\lambda_4 = 2$	Decay rate for prior variance on lags
$\lambda_5 = 0.001$	Relative precision for intercept

Instead of working with covariance matrices, it is more efficient to work with precision matrices. Define:

$$\underline{P}_i^\beta = (\underline{V}_i^\beta)^{-1}, \quad \bar{P}_i^\beta = (\bar{V}_i^\beta)^{-1}.$$

The updating equations for the posterior take the form

$$\begin{aligned} \bar{P}_i^\beta &= \underline{P}_i^\beta + Z_i' Z_i \\ \bar{\beta}_i &= (\bar{P}_i^\beta)^{-1} (\underline{P}_i^\beta \underline{\beta}_i + Z_i' W_i) \\ \bar{\nu}_i &= \underline{\nu}_i + T/2 \\ \bar{S}_i &= \underline{S}_i + \frac{1}{2} (W_i' W_i + \underline{\beta}_i' \underline{P}_i^\beta \underline{\beta}_i - \bar{\beta}_i' \bar{P}_i^\beta \bar{\beta}_i). \end{aligned}$$

When using the JK instruments, we set the coefficients on the lags and the intercept equal to zero: $B_i = 0$ for $i = 1, 2$. Thus, $\beta_1 = 0_{(np+1) \times 1}$ and $\beta_2 = [A_{2,1}, 0_{1 \times (np+1)}]'$. The updating equations for the posterior change as follows. For $i = 1$:

$$\begin{aligned} \bar{P}_1^\beta &= N/A \\ \bar{\beta}_1 &= 0_{(np+1) \times 1} \\ \bar{\nu}_1 &= \underline{\nu}_i + T/2 \\ \bar{S}_1 &= \underline{S}_1 + \frac{1}{2} W_1' W_1. \end{aligned}$$

For $i = 2$ we are regressing W_{2t} on the single regressor W_{1t} . Partition $\beta_2 = [\beta_{1,2}, 0_{1 \times (np+1)}]'$ and denote the precision associated with the first element of the β_2 vector by $[P_1^\beta]_{11}$. Then,

we can write the updating equations as

$$\begin{aligned}
[\bar{P}_2^\beta]_{11} &= [\underline{P}_2^\beta]_{11} + W_1' W_1 \\
\bar{\beta}_{2,1} &= ([\bar{P}_2^\beta]_{11})^{-1} ([\underline{P}_2^\beta]_{11} \underline{\beta}_{2,1} + W_1' W_2) \\
\bar{\nu}_2 &= \underline{\nu}_2 + T/2 \\
\bar{S}_2 &= \underline{S}_2 + \frac{1}{2} (W_2' W_2 + [\underline{P}_2^\beta]_{11} \underline{\beta}_{2,1}^2 - [\bar{P}_2^\beta]_{11} \bar{\beta}_{2,1}^2).
\end{aligned}$$

The MDD can be computed analytically as follows:

$$\begin{aligned}
\ln p(W) &= -\frac{Tn}{2} \ln(2\pi) + \sum_{i=1}^n \left[\frac{1}{2} (\ln |\underline{P}_i^\beta| - \ln |\bar{P}_i^\beta|) \right. \\
&\quad \left. + \underline{\nu}_i \ln |\underline{S}_i| - \bar{\nu}_i \ln |\bar{S}_i| - \ln \Gamma(\underline{\nu}_i) + \ln \Gamma(\bar{\nu}_i) \right],
\end{aligned} \tag{A.26}$$

with the understanding that for $i = 1$ the expression $\ln |\underline{P}_i^\beta| - \ln |\bar{P}_i^\beta| = 0$ and for $i = 2$ it gets replaced by $\ln [\underline{P}_i^\beta]_{11} - \ln [\bar{P}_i^\beta]_{11}$.

B.3 Separating Conventional and Informational Policy Shocks

It is assumed that a contractionary monetary policy shock generates an interest rate increase and a drop in stock prices, whereas a positive information shock is associated with an increase in both interest rates and stock prices. As in JK, to separate monetary policy shocks from information shocks one has to introduce a 2×2 orthogonal matrix Ω that links the orthogonalized instrument innovations to the structural shocks. Our Bayesian analysis requires a prior distribution $p(\Omega|\beta, D)$. Because Ω does not affect the likelihood function, conditional on the reduced-form VAR parameters, its prior does not get updated. We use a prior for Ω that is uniform over the identified set defined by the sign restrictions, conditional on the reduced-form parameters.

C State-Space Representation for Mixed-Frequency Analysis

C.1 Companion Form

To derive the updating equations for the filter and simulation smoother we express the state-transition equation in companion form. We illustrate the companion form notation for $p = 2$.

The generalization is straightforward. We define

$$W'_t = [Y'_t, Y'_{t-1}, \alpha'_t, \alpha'_{t-1}] \quad (\text{A.27})$$

and partition W'_t into

$$W'_t = [Z'_t, s'_t], \quad Z'_t = [Y'_t, Y'_{t-1}], \quad s_t = [\alpha_t, \alpha_{t-1}]'.$$

The companion-form law of motion for w_t can be written as

$$\begin{bmatrix} Y_t \\ Y_{t-1} \\ \alpha_t \\ \alpha_{t-1} \end{bmatrix} = \begin{bmatrix} \Phi_{1,yy} & \Phi_{2,yy} & \Phi_{1,y\alpha} & \Phi_{2,y\alpha} \\ I_y & 0 & 0 & 0 \\ \Phi_{1,\alpha y} & \Phi_{2,\alpha y} & \Phi_{1,\alpha\alpha} & \Phi_{2,\alpha\alpha} \\ 0 & 0 & I_\alpha & 0 \end{bmatrix} \begin{bmatrix} Y_{t-1} \\ Y_{t-2} \\ \alpha_{t-1} \\ \alpha_{t-2} \end{bmatrix} + \begin{bmatrix} I_y & 0 \\ 0 & 0 \\ 0 & I_\alpha \\ 0 & 0 \end{bmatrix} \begin{bmatrix} u_{y,t} \\ u_{\alpha,t} \end{bmatrix}, \quad (\text{A.28})$$

or, more compactly, as

$$W_t = \Psi W_{t-1} + M u_t. \quad (\text{A.29})$$

We further partition the selection matrix M into

$$M = [M_y \ M_\alpha]$$

such that

$$M'_y W_t = Y_t, \quad M'_\alpha W_t = \alpha_t.$$

In some instances, we need to extract the subvectors Z_t and s_t from W_t , which is done by the selection matrices

$$\Xi'_z = [I_z \ 0], \quad Z_t = \Xi'_z W_t, \quad \Xi'_s = [0 \ I_s], \quad s_t = \Xi'_s W_t.$$

For the forward iterations of the Kalman filter, it is useful to separate s_t and Z_t . To that end, define

$$\Psi_{ss} = \Xi'_s \Psi \Xi_s = \begin{bmatrix} \Phi_{1,\alpha\alpha} & \Phi_{2,\alpha\alpha} \\ I_\alpha & 0 \end{bmatrix}, \quad \Psi_{sz} = \Xi'_s \Psi \Xi_z = \begin{bmatrix} \Phi_{1,\alpha y} & \Phi_{2,\alpha y} \\ 0 & 0 \end{bmatrix}.$$

Moreover, we define

$$M_{s\alpha} = \Xi'_s M_\alpha = \begin{bmatrix} I_\alpha \\ 0 \end{bmatrix}, \quad \Phi_{yz} = [\Phi_{1,yy} \ \Phi_{2,yy}], \quad \Phi_{ys} = [\Phi_{1,y\alpha} \ \Phi_{2,y\alpha}].$$

Using this notation, the measurement equation can be written as

$$\hat{\alpha}_t = M'_{s\alpha} s_t + \eta_t, \quad \eta_t \sim N(0, V_t). \quad (\text{A.30})$$

C.2 Forward Filtering

The forward filtering iterations are obtained from a modified Kalman filter that recognizes that y_t is directly observable. Thus, only s_t is a latent state variable.

Recursive Assumption. We start from the assumption that

$$s_{t-1}|(Y_{1:t-1}, \hat{\alpha}_{1:t-1}) \sim \mathcal{N}(s_{t-1|t-1}, P_{t-1|t-1}(ss)). \quad (\text{A.31})$$

We use the notation $P_{t-1|t-1}(ss)$ to indicate that one can define a larger matrix, conforming with w_t , which is partitioned as follows:

$$P_{t-1|t-1} = \begin{bmatrix} P_{t-1|t-1}(zz) & P_{t-1|t-1}(zs) \\ P_{t-1|t-1}(zs) & P_{t-1|t-1}(ss) \end{bmatrix},$$

with the understanding that

$$P_{t-1|t-1}(zz) = 0, \quad P_{t-1|t-1}(zs) = 0, \quad P_{t-1|t-1}(sz) = 0.$$

Initialization. Just in a regular VAR analysis, we will condition the likelihood function on observations to initialize the lags associated with the determination of (Z_1, s_1) . Note that the initial lags of Y_t are directly observed under suitable definition of the sample period and hence we know Z_0 . To initialize the latent state we set

$$s_{0|0} = \begin{bmatrix} \hat{\alpha}_0 \\ \vdots \\ \hat{\alpha}_{-p+1} \end{bmatrix}, \quad P_{0|0}(ss) = \text{diag}[V_0, \dots, V_{-p+1}].$$

Thus, we assume that

$$s_0|(Y_{-p+1:0}, \hat{\alpha}_{-p+1:0}) \sim \mathcal{N}(s_{0|0}, P_{0|0}).$$

For the smoother below, it will become important to properly account for the initialization because we will also need draws of $s_0 = \alpha_{-p+1:0}$ to set up the Gibbs sampler.

Forecasting. We begin by forecasting (Y_t, s_t) jointly, using (A.28). Let

$$\begin{bmatrix} Y_t \\ s_t \end{bmatrix} \Big| (Y_{-p+1:t-1}, \hat{\alpha}_{-p+1:t-1}) \sim \mathcal{N} \left(\begin{bmatrix} y_{t|t-1} \\ s_{t|t-1} \end{bmatrix}, \begin{bmatrix} P_{t|t-1}(yy) & P_{t|t-1}(ys) \\ P_{t|t-1}(sy) & P_{t|t-1}(ss) \end{bmatrix} \right), \quad (\text{A.32})$$

where

$$\begin{aligned}
Y_{t|t-1} &= \Phi_{yz}Z_{t-1} + \Phi_{ys}s_{t-1|t-1} \\
s_{t|t-1} &= \Psi_{sz}Z_{t-1} + \Psi_{ss}s_{t-1|t-1} \\
P_{t|t-1}(yy) &= \Phi_{ys}P_{t-1|t-1}(ss)\Phi'_{ys} + \Sigma_{yy} \\
P_{t|t-1}(sy) &= \Psi_{ss}P_{t-1|t-1}(ss)\Phi'_{ys} + M_{s\alpha}\Sigma_{\alpha y} \\
P_{t|t-1}(ss) &= \Psi_{ss}P_{t-1|t-1}(ss)\Psi'_{ss} + M_{s\alpha}\Sigma_{\alpha\alpha}M'_{s\alpha}.
\end{aligned}$$

We can now factorize the joint distribution of (Y_t, s_t) into a marginal distribution of Y_t and a conditional distribution of $s_t|Y_t$:

$$\begin{aligned}
Y_t|(Y_{-p+1:t-1}, \hat{\alpha}_{-p+1:t-1}) &\sim \mathcal{N}(Y_{t|t-1}, P_{t|t-1}(yy)) \\
s_t|(Y_t, Y_{-p+1:t-1}, \hat{\alpha}_{-p+1:t-1}) &\sim \mathcal{N}(s_{t|y,t-1}, P_{t|y,t-1}(ss)),
\end{aligned} \tag{A.33}$$

where

$$\begin{aligned}
s_{t|y,t-1} &= s_{t|t-1} + P_{t|t-1}(sy)[P_{t|t-1}(yy)]^{-1}(Y_t - Y_{t|t-1}) \\
P_{t|y,t-1}(ss) &= P_{t|t-1}(ss) - P_{t|t-1}(sy)[P_{t|t-1}(yy)]^{-1}P_{t|t-1}(ys).
\end{aligned}$$

Finally, the distribution of $\hat{\alpha}_t$ conditional on y_t and $t-1$ information is

$$\hat{\alpha}_t|(Y_t, Y_{-p+1:t-1}, \hat{\alpha}_{-p+1:t-1}) \sim N(\hat{\alpha}_{t|y,t-1}, F_{t|y,t-1}), \tag{A.34}$$

where

$$\begin{aligned}
\hat{\alpha}_{t|y,t-1} &= M'_{s\alpha}s_{t|y,t-1} \\
F_{t|y,t-1} &= M'_{s\alpha}P_{t|y,t-1}(ss)M_{s\alpha} + V_t.
\end{aligned}$$

Updating. The updating step in the Kalman filter is done conditional on Y_t . Generically,

$$\begin{aligned}
&p(s_t|\hat{\alpha}_t, Y_t, Y_{-p+1:t-1}, \hat{\alpha}_{-p+1:t-1}) \\
&\propto p(\hat{\alpha}_t|s_t, Y_t, Y_{-p+1:t-1}, \hat{\alpha}_{-p+1:t-1})p(s_t|y_t, Y_{-p+1:t-1}, \hat{\alpha}_{-p+1:t-1}).
\end{aligned} \tag{A.35}$$

We can use the standard updating formulas, substituting in the means and variances that are conditional on Y_t :

$$\begin{aligned}
s_{t|t} &= s_{t|y,t-1} + P_{t|y,t-1}(ss)M_{s\alpha}F_{t|y,t-1}^{-1}(\hat{\alpha}_t - \hat{\alpha}_{t|y,t-1}) \\
P_{t|t}(ss) &= P_{t|y,t-1}(ss) - P_{t|y,t-1}(ss)M_{s\alpha}F_{t|y,t-1}^{-1}M'_{s\alpha}P_{t|y,t-1}(ss).
\end{aligned} \tag{A.36}$$

C.3 Simulation Smoothing

The goal is to generate draws from the distribution of $\alpha_{-p+1:T}$ given $(Y_{-p+1:T}, \hat{\alpha}_{-p+1:T})$. Note that the filter in its final step generates the distribution $s_T | (Y_{-p+1:T}, \hat{\alpha}_{-p+1:T})$. Because of the companion form, a draw s_T^i determines $\alpha_T^i, \dots, \alpha_{T-p+1}^i$. Thus, the distribution $s_{T-1} | (s_T, Y_{-p+1:T}, \hat{\alpha}_{-p+1:T})$ is degenerate. In the implementation of the smoother, we will draw blocks as follows:

$$s_T, \quad s_{T-p} | s_T, \quad s_{T-2p} | (s_{T-p}, s_T), \quad s_{T-3p} | (s_{T-2p}, s_{T-p}, s_T), \quad \dots$$

Because $s_{T-jp} = [\alpha'_{T-jp}, \dots, \alpha'_{T-(j+1)p+1}]'$, this approach generates the entire sequence $\alpha_{-p+1:T}$. However, special care needs to be given to $\alpha_{-p+1:0}$ and the fact that T may not be divisible by p .

Preliminaries. Assuming for now that T is a multiple of p and that $T = bp$, where b is the number of blocks, the simulation smoother relies on the factorization

$$p(Y_{-p+1:T}, \alpha_{-p+1:T} | Y_{-p+1:0}, \alpha_{-p+1:0}) = \prod_{j=1}^b p(Y_{(j-1)p+1:j p}, \alpha_{(j-1)p+1:j p} | Y_{-p+1:(j-1)p}, \alpha_{-p+1:(j-1)p}),$$

where, because of the VAR(p) structure of the state-transition equation,

$$\begin{aligned} & p(Y_{(j-1)p+1:j p}, \alpha_{(j-1)p+1:j p} | Y_{-p+1:(j-1)p}, \alpha_{-p+1:(j-1)p}) \\ &= p(Y_{(j-1)p+1:j p}, \alpha_{(j-1)p+1:j p} | Y_{(j-2)p+1:(j-1)p}, \alpha_{(j-2)p+1:(j-1)p}). \end{aligned}$$

The conditional distribution on the right-hand side, can be obtained by iterating the companion form (A.29) p periods forward:

$$W_{t+p} = \Psi^p W_t + M \sum_{h=0}^{p-1} \Psi^h u_{t+p-h}.$$

Thus,

$$W_{t+p} | W_t \sim \mathcal{N} \left(\Psi^p W_t, \sum_{h=0}^{p-1} \Psi^h M \Sigma M' \Psi^{h'} \right). \quad (\text{A.37})$$

Now, set $t = (j-1)p$ and note that

$$W_{(j-1)p} = [Y_{(j-1)p}, \alpha_{(j-1)p}, \dots, Y_{(j-2)p+1}, \alpha_{(j-2)p+1}],$$

as required.

Generic Smoother. We first modify the generic derivation of the smoother to account for the presence of both $Y_{-p+1:T}$ and $\hat{\alpha}_{-p+1:T}$ in the conditioning set and the block sampling. As

before, it is convenient for now to assume that $T = bp$ and the time index t shifts in steps of p periods, that is, $t = (j - 1)p$, where $j = 0, \dots, b$. Consider the following factorization

$$\begin{aligned}
& p(\alpha_{t-p+1:t}, \alpha_{t+1:T}, Y_{-p+1:T}, \hat{\alpha}_{-p+1:T}) \\
&= \int p(\alpha_{-p+1:T}, Y_{-p+1:T}, \hat{\alpha}_{-p+1:T}) d\alpha_{-p+1:t-p} \\
&= \int p(\alpha_{-p+1:t}, Y_{-p+1:t}, \hat{\alpha}_{-p+1:t}) \\
&\quad \times \left(\prod_{\tau=t+1}^T p(Y_\tau | Y_{\tau-p:\tau-1}, \alpha_{\tau-p:\tau-1}) p(\alpha_\tau | Y_\tau, Y_{\tau-p:\tau-1}, \alpha_{\tau-p:\tau-1}) p(\hat{\alpha}_\tau | \alpha_\tau) \right) d\alpha_{-p+1:t-p} \\
&= p(\alpha_{t-p+1:t}, Y_{-p+1:t}, \hat{\alpha}_{-p+1:t}) \\
&\quad \times \prod_{\tau=t+1}^{t+p} p(Y_\tau | Y_{\tau-p:\tau-1}, \alpha_{\tau-p:\tau-1}) p(\alpha_\tau | Y_\tau, Y_{\tau-p:\tau-1}, \alpha_{\tau-p:\tau-1}) p(\hat{\alpha}_\tau | \alpha_\tau) \\
&\quad \times \text{terms without } \alpha_{t-p+1:t}.
\end{aligned}$$

Maintaining the assumption that $T = bp$ and $t = (j - 1)p$, we can deduce

$$\begin{aligned}
& p(s_t | s_{t+p}, s_{t+2p}, \dots, s_T, Y_{-p+1:T}, \hat{\alpha}_{-p+1:T}) \tag{A.38} \\
&= p(\alpha_{t-p+1:t} | \alpha_{t+1:T}, Y_{-p+1:T}, \hat{\alpha}_{-p+1:T}) \\
&\propto p(\alpha_{t-p+1:t}, \alpha_{t+1:T}, Y_{-p+1:T}, \hat{\alpha}_{-p+1:T}) \\
&\propto p(\alpha_{t-p+1:t}, Y_{-p+1:t}, \hat{\alpha}_{-p+1:t}) p(Y_{t+1:t+p}, \alpha_{t+1:t+p} | Y_{t-p+1:t}, \alpha_{t-p+1:t}) \\
&= p(s_t, Y_{-p+1:t}, \hat{\alpha}_{-p+1:t}) p(Z_{t+p}, s_{t+p} | Z_t, s_t).
\end{aligned}$$

The first proportionality follows from Bayes Theorem. The second proportionality follows from dropping the factor $p(Y_{t+p+1:T}, \alpha_{t+p+1:T} | Y_{t+1:t+p}, \alpha_{t+1:t+p})$ because it does not depend on $\alpha_{t-p+1:t}$. The last equality uses $s_t = [\alpha'_{t-p+1}, \dots, \alpha'_t]'$ and $Z_t = [Y'_{t-p+1}, \dots, Y'_t]'$. Because

$$p(s_t, Y_{-p+1:t}, \hat{\alpha}_{-p+1:t}) = p(s_t | Y_{-p+1:t}, \hat{\alpha}_{-p+1:t}) p(Y_{-p+1:t}, \hat{\alpha}_{-p+1:t}),$$

it follows that

$$p(s_t | s_{t+p}, s_{t+2p}, \dots, s_T, Y_{-p+1:T}, \hat{\alpha}_{-p+1:T}) \propto p(s_t | Y_{-p+1:t}, \hat{\alpha}_{-p+1:t}) p(Z_{t+p}, s_{t+p} | Z_t, s_t). \tag{A.39}$$

Smoothing Formulas for the Linear Gaussian Model. We previously established that

$$s_t | (Y_{1:t}, \hat{\alpha}_{1:t}) \sim \mathcal{N}(s_t | t, P_t | t(ss)).$$

As previously shown in (A.37), iterating the companion form forward for p periods yields

$$(Z_{t+p}, s_{t+p}) | (Z_t, s_t) \sim \mathcal{N} \left(\Psi^p W_t, \sum_{h=0}^{p-1} \Psi^h M \Sigma M' \Psi^{h'} \right). \tag{A.40}$$

Let

$$W_{t|t} = [Z'_t, s'_{t|t}]'$$

and define

$$P_{t+p|t}(ww) = \Psi^p \Xi_s P_{t|t}(ss) \Xi'_s \Psi^p + \sum_{h=0}^{p-1} \Psi^h M \Sigma M' \Psi^{h'}. \quad (\text{A.41})$$

The joint distribution of (s_t, W_{t+p}) is given by

$$\begin{bmatrix} s_t \\ W_{t+p} \end{bmatrix} \Big| (\cdot) \sim \mathcal{N} \left(\begin{bmatrix} s_{t|t} \\ \Psi^p W_{t|t} \end{bmatrix}, \begin{bmatrix} P_{t|t}(ss) & P_{t|t}(ss) \Xi'_s \Psi^{p'} \\ \Psi^p \Xi_s P_{t|t}(ss) & P_{t+p|t}(ww) \end{bmatrix} \right). \quad (\text{A.42})$$

Then, we sample s_t from

$$s_t \sim \mathcal{N}(s_{t|t+p}^i, P_{t|t+p}(ss)), \quad (\text{A.43})$$

where

$$s_{t|t+p} = s_{t|t} + P_{t|t}(ss) \Xi'_s \Psi^{p'} P_{t+p|t}^{-1}(ww) \left(\begin{bmatrix} Z_{t+p} \\ s_{t+p} \end{bmatrix} - \Psi^p W_{t|t} \right) \quad (\text{A.44})$$

$$P_{t|t+p}(ss) = P_{t|t}(ss) - P_{t|t}(ss) \Xi'_s \Psi^{p'} P_{t+p|t}^{-1}(ww) \Psi^p \Xi_s P_{t|t}(ss). \quad (\text{A.45})$$

T is not a multiple of p . Consider the following example. Suppose that $T = 10$ and $p = 3$.

In this case we can use the formulas to generate

$$\begin{aligned} p(s_{10}|Y_{-2:10}, \hat{\alpha}_{-2:10}), \quad p(s_7|s_{10}, Y_{-2:10}, \hat{\alpha}_{-2:10}), \\ p(s_4|s_7, s_{10}, Y_{-2:10}, \hat{\alpha}_{-2:10}), \quad p(s_1|s_4, s_7, s_{10}, Y_{-2:10}, \hat{\alpha}_{-2:10}). \end{aligned}$$

The last step gives us

$$p(\alpha_{-1:1}|\alpha_{2:10}, Y_{-2:10}, \hat{\alpha}_{-2:10}),$$

but we still lack

$$p(\alpha_{-2}|\alpha_{-1:10}, Y_{-2:10}, \hat{\alpha}_{-2:10}).$$

In order to generate the draws needed to initialize the vector autoregressive law of motion of the state transition equation, we take the following short-cut: (i) discard $\alpha_{-1:0}$, (ii) redraw s_0 from (A.43) by evaluating the formulas $s_{t|t+p}$ and $P_{t|t+p}(ss)$ for $t = 0$.

D Data Used in the Empirical Analysis

Aggregate Data. Following Jarocinski and Karadi (2020b) we use six monthly macroeconomic variables in the empirical model: (i) the monthly average of the one-year constant-maturity Treasury yield serves as the monetary policy indicator. (ii) The monthly average

of the S&P 500 stock price index in log levels. (iii,iv) Real GDP and GDP deflator in log levels interpolated to monthly frequency based on Stock and Watson (2010). (v) The excess bond premium (EBP) as indicator of financial conditions. (vi) An aggregate employment rate constructed from the micro data (see below). The functional VAR with micro-level consumption data includes an additional aggregate variable: real personal consumption expenditures per capita in log levels. We use Personal Consumption Expenditures (*PCE*, U.S. Bureau of Economic Analysis (2021)) and divide it by population level (*CNP16OV*, U.S. Bureau of Labor Statistics (2021)) to get per capita values. Then we use GDP deflator to get the real values. Data are obtained from Jarocinski and Karadi (2020a) and Federal Reserve Bank of St. Louis (2022).

Micro-Level Earnings Data. The CPS raw data are downloaded from <http://www.nber.org/data/cps-basic.html>. Reference: U.S. Census Bureau and U.S. Bureau of Labor Statistics (2022).

The raw data files are converted into STATA using the do-files available at:

<http://www.nber.org/data/cps-basic-progs.html>.

We use the series PREXPLF (“Experienced Labor Force Employment”), which is the same as in the raw data, and the series PRERNWA (“Weekly Earnings”), which is constructed as PEHRUSL1 (“Hours Per Week at One’s Main Job”) times PRHERNAL (“Hourly Earnings”) for hourly workers, and given by PRWERNAL for weekly workers. STATA dictionary files are available at: <http://www.nber.org/data/progs/cps-basic/>

We pre-process the cross-sectional data as follows. We drop individuals if (i) the employment indicator is not available; and (ii) if they are coded as “employed” but the weekly earnings are missing. In addition, we re-code individuals with non-zero earnings as employed and set earnings to zero for individuals that are coded as not employed. A CPS-based unemployment rate is computed as the fraction of individuals that are coded as not employed. By construction this is one minus the fraction of individuals with non-zero weekly earnings, which is used to normalize the cross-sectional density of earnings. It turns out that the CPS-based unemployment rate tracks the aggregate unemployment rate (*UNRATE*, National Bureau of Economic Research (????)) very closely. The levels of the two series are very similar, but the CPS unemployment rate exhibits additional high-frequency fluctuations, possibly due to seasonals that have been removed from the aggregate unemployment rate.

Below we also report results based on quarterly labor earnings from the CPS. To generate these results, we use the quarterly version of the CPS labor earnings data constructed in

Chang, Chen, and Schorfheide (2024).

Micro-Level Hours Worked Data. We use the CPS series PEHRUSL1, which was previously used to construct the micro-level earnings. No further transformations are applied.

Micro-Level Consumption Data. We use public use microdata (PUMD) from the CEX conducted by the U.S. Bureau of Labor Statistics (2017). Although interviews about consumption expenditures are conducted at monthly frequency, we aggregate to quarterly frequency to smooth out unit-level expenditures and reduce reporting errors. Each consumption unit (CU) provides expenditure data for three consecutive months. If a CU was surveyed in 1990Q1, consumption responses could refer to one of three possible combination of months: (i) 1989:M10, 1989:M11, 1989:M12; (ii) 1989:M11, 1989:M12, 1990:M1; or (iii) 1989:M12, 1990:M1, 1990:M2. To convert this information into quarterly expenditure data, we add the three monthly values and assign the sum to the quarter that covers at least two of the months for which responses were obtained. We divide the expenditures by the number of individuals with age 16 and over belonging to the consumption unit (i.e. household or family) to obtain per capita expenditures. Because CEX per capita expenditures capture less than 50% of NIPA per capita expenditures we rescale them as follows. Let C_t be aggregate per capita consumption from NIPA. We calculate $\frac{1}{T} \sum_{t=1}^T \text{median}(c_{1t}, \dots, c_{Nt})/C_t \approx 0.46$. We then define $z_{it} = c_{it}/(0.46 \cdot C_t)$. Thus, if $z_{it} = 1$ the individual approximately consumes at the level of aggregate consumption per capita. Finally, as for the earnings data, we apply the inverse hyperbolic sine transformation to obtain x_{it} .

Micro-Level Financial Data. We are using the Distributional National Accounts micro-files provided by Piketty, Saez, and Zucman (2018b). Reference: Piketty, Saez, and Zucman (2018a). We use a broad measure of financial income defined as *fkinc*: personal factor capital income = *fkhou* (housing asset income) + *fkequ* (equity asset income, S corporations) + *fkfix* (interest income) + *fkbus* (business asset income) + *fkpen* (pension and insurance asset income) + *fkdeb* (interest payments) + *fkprk* (sales and excise taxes allocated to capital) + *fksubk* (subsidies allocated to capital). We transform the financial income data by first dividing by (1/3 of GDP) and then applying the inverse hyperbolic sine transform.

Inverse Hyperbolic Sign Transformation. We transform the micro data (labor or financial income to GDP ratios, consumption relative to aggregate per capita consumption), denoted by z below, using the inverse hyperbolic sine transformation, which is given by

$$x = g(z|\theta) = \frac{\ln(\theta z + (\theta^2 z^2 + 1)^{1/2})}{\theta} = \frac{\sinh^{-1}(\theta z)}{\theta} \quad (\text{A.46})$$

with $\theta = 1$. Note that $g(0|\theta) = 0$ and $g^{(1)}(0|\theta) = 1$, that is, for small values of z the transformation is approximately linear. For large values of z the transformation is logarithmic:

$$g(z|\theta) \approx \frac{1}{\theta} \ln(2\theta z) = \frac{1}{\theta} \ln(2\theta) + \frac{1}{\theta} \ln(z).$$

The inverse of the transformation takes the form

$$z = g^{-1}(x|\theta) = \frac{1}{\theta} \sinh(\theta x) = \frac{1}{2\theta} (e^{\theta x} - e^{-\theta x}).$$

Most of the calculations in the paper are based on $p_x(x)$. But in some instances, it is desirable to report for $p_z(z)$. From a change of variables (omitting the θ), we get

$$p_z(z) = p_x(g(z)) |g'(z)|,$$

where

$$g'(z) = \frac{1 + \frac{\theta z}{(\theta^2 z^2 + 1)^{1/2}}}{\theta z + (\theta^2 z^2 + 1)^{1/2}} = \frac{1}{(\theta^2 z^2 + 1)^{1/2}}.$$

Whenever we do convert the estimated densities back from z to x , we recycle the density evaluations at x_j . Thus, we evaluate $p_z(z)$ for grid points $z_j = g^{-1}(x_j)$, which leads to

$$p_z(z_j) = p_x(x_j) |g'(g^{-1}(x_j))|,$$

where

$$|g'(g^{-1}(x_j))| = \frac{1}{\left(\frac{1}{4}(e^{\theta x_j} - e^{-\theta x_j})^2 + 1\right)^{1/2}} = \frac{2}{(e^{2\theta x_j} + e^{-2\theta x_j} + 2e^{2\theta x_j} e^{-2\theta x_j})^{1/2}} = \frac{2}{e^{\theta x_j} + e^{-\theta x_j}}.$$

Narrative Evidence. To search the FOMC minutes for evidence on concerns about inequality influencing the interest rate decision, we proceeded as follows: For the period 1990–2016, we used the Federal Reserve Board’s historical archive to access FOMC meeting minutes in HTML format. Specifically, we relied on the following webpage to locate the minutes for each meeting: https://www.federalreserve.gov/monetarypolicy/fomc_historical_year.htm. For example, the minutes for the January 26–27, 2016 meeting are available at <https://www.federalreserve.gov/monetarypolicy/fomcminutes20160127.htm>.

For each meeting date, we provided the corresponding link to ChatGPT and asked it to search the minutes for the following keywords: “inequality”, “distribution”, “distributional”, “uneven”, “disparities”, and “lower-income”. We did not find any instances of these terms appearing in the minutes during this period.

Given the absence of these keywords in the FOMC minutes, ChatGPT suggested expanding the search to other sources, such as speeches or interviews by FOMC officials. Following this approach, we identified speeches by Chair Janet Yellen in 2014 that explicitly discuss distributional and inequality-related issues.

Finally, although it falls outside the sample period, we found the following relevant passage in the 2020:M6 FOMC minutes: “The Committee noted that the burdens of the COVID-19 downturn were not falling equally across the population, with job losses being especially severe for lower-wage workers, women, African Americans, and Hispanics.”

While the minutes do not use the specific terms “inequality”, “distributional”, or “uneven”, this language clearly reflects an awareness of heterogeneous labor market impacts across groups.

As a validity check, we also asked ChatGPT to search for the keyword “inflation” in the same set of minutes. In this case, it successfully identified multiple occurrences, confirming that the keyword-search procedure itself is functioning as intended.

E Knot Placement, Lag length, and Hyperparameters

- Figure 1: the results for the aggregate VAR are generated from the following specification: $\hat{p} = 4$, $\hat{\lambda}_1 = 148.4$, and $\hat{\lambda}_2 = 1$.
- Figure 5 (functional VAR with monthly earnings) is generated with the following settings:

- Knot placement (percentiles of distribution of pooled across time periods micro-level data):

$$K = 4 \quad : \quad 0.25, 0.50, 0.75$$

$$K = 6 \quad : \quad 0.10, 0.25, 0.50, 0.75, 0.90$$

$$K = 8 \quad : \quad 0.05, 0.10, 0.25, 0.50, 0.75, 0.90, 0.95$$

$$K = 10 \quad : \quad 0.01, 0.025, 0.05, 0.10, 0.25, 0.50, 0.75, 0.90, 0.95$$

- Splines: $x \in [0, 3]$, cubic-cubic-linear constructed from the right.
- Model specifications for MDD selection: We consider four values of K : 4, 6, 8, and 10. For λ_1 and λ_2 , we use 31 equally spaced values of $\ln \lambda_j$ between -10 and 20, while λ_3 is fixed at 1. The lag length p takes four values: 1, 2, 3, and 4.
- Model selection: $K = 10$, $p = 1$, $\lambda_1 = 54.6$, $\lambda_2 = 54.6$, $\lambda_3 = 1$.

- Figure 7 (functional VAR with monthly hours worked) is generated with the following settings:

- Knot placement (percentiles of distribution of pooled across time periods micro-level data):

$$K = 4 \quad : \quad 0.25, 0.50, 0.75$$

$$K = 6 \quad : \quad 0.10, 0.25, 0.50, 0.75, 0.90$$

$$K = 8 \quad : \quad 0.05, 0.10, 0.25, 0.50, 0.75, 0.90, 0.95$$

$$K = 10 \quad : \quad 0.01, 0.025, 0.05, 0.10, 0.25, 0.50, 0.75, 0.90, 0.95$$

- Splines: $x \in [1.05, 155.8]$, cubic-cubic-linear constructed from the right.
- Model selection: $K = 10$, $p = 1$, $\lambda_1 = 54.6$, $\lambda_2 = 20.1$, $\lambda_3 = 1$.

- Figure 9 (functional VAR with quarterly total consumption) is generated with the following settings:

- Knot placement (percentiles of distribution of pooled across time periods micro-level data):

$$K = 4 \quad : \quad 0.25, 0.50, 0.75$$

$$K = 6 \quad : \quad 0.10, 0.25, 0.50, 0.75, 0.90$$

$$K = 8 \quad : \quad 0.05, 0.10, 0.25, 0.50, 0.75, 0.90, 0.95$$

$$K = 10 \quad : \quad 0.01, 0.025, 0.05, 0.10, 0.25, 0.50, 0.75, 0.90, 0.95$$

- Splines: $x \in [0, 3]$, cubic-cubic-linear constructed from the right.
- Model selection: $K = 6$, $p = 1$, $\lambda_1 = 403.4$, $\lambda_2 = 0.37$, $\lambda_3 = 1$.
- Figure 10 (functional VARs with quarterly consumption components) is generated with the following settings:
 - Model selection (Total): $K = 6$, $p = 1$, $\lambda_1 = 403.4$, $\lambda_2 = 0.37$, $\lambda_3 = 1$.
 - Model selection (Durables): $K = 8$, $p = 1$, $\lambda_1 = 54.6$, $\lambda_2 = 20.1$, $\lambda_3 = 1$.
 - Model selection (Nondurables): $K = 10$, $p = 1$, $\lambda_1 = 54.6$, $\lambda_2 = 7.4$, $\lambda_3 = 1$.
 - Model selection (Services): $K = 4$, $p = 1$, $\lambda_1 = 403.4$, $\lambda_2 = 0.37$, $\lambda_3 = 1$.
- The functional VAR with stacked earnings and consumption densities is done with the following hyperparameters (based on MDD): $\lambda_1 = 54.6$, $\lambda_2 = 54.6$, $K_E = 8$, $K_C = 6$.
- Figure 13 (mixed-frequency functional VAR with quarterly aggregate variables and annual financial income density) is generated with the following settings:
 - Knot placement (percentiles of distribution of pooled across time periods micro-level data):

$$K = 11 \quad : \quad 0.005, 0.01, 0.05, 0.15, 0.2, 0.3, 0.75, 0.9, 0.95, 0.99.$$

- Splines: $x \in [-5, 10]$, linear-cubic-linear constructed from the left.
- Model choice: $K = 11$, $p = 1$, $\lambda_1 = 148.4$, $\lambda_2 = 2.71$, $\lambda_3 = 1$.

F Additional Empirical Results

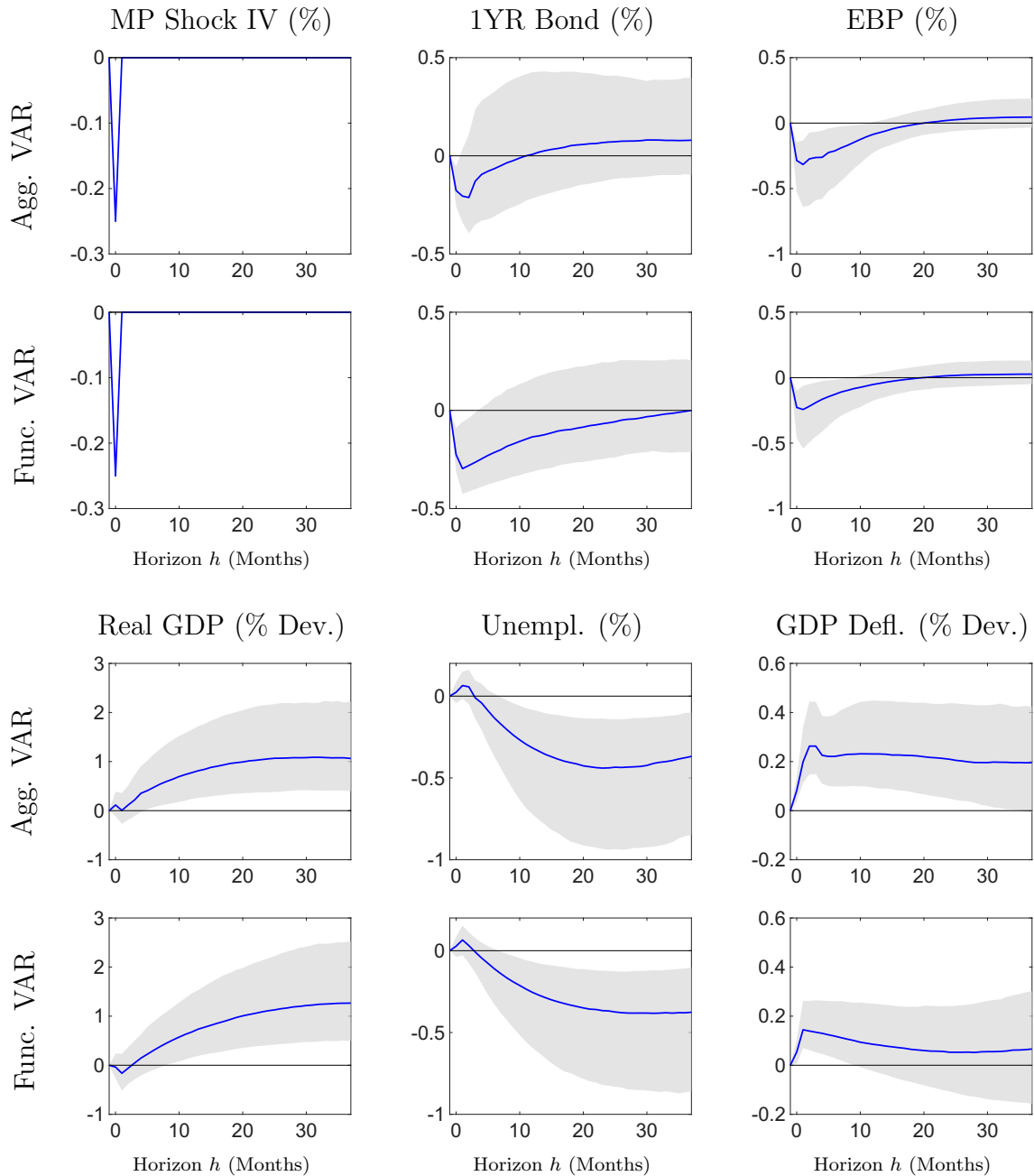
F.1 IRFs of Aggregate Variables

Figure A-1 compares the responses of the aggregate variables estimated from a functional VAR that includes the labor earnings distribution and a standard VAR that only includes the macroeconomic aggregates. The responses are generally very similar. The presence of the cross-sectional variables has no substantial effect on the impulse response inference. The main difference is that the GDP deflator response in the functional VAR is slightly weaker. The width of the bands depends on the number of lags and the hyperparameters selected by the MDD. The results for the aggregate VAR are based on $p = 4$, whereas the functional VAR (earnings) uses $p = 1$.

Figure A-2 plots the response of aggregate variables to a monetary policy shock for other model specifications discussed in the main paper. Note that the functional VAR for hours worked (plots in the first row of the figure) is estimated based on monthly data, which is why the horizon ranges from zero to 36 months (3 years). The remaining functional VARs are estimated based on quarterly data and the IRF horizons correspond to quarters, ranging from zero to 12 (3 years).

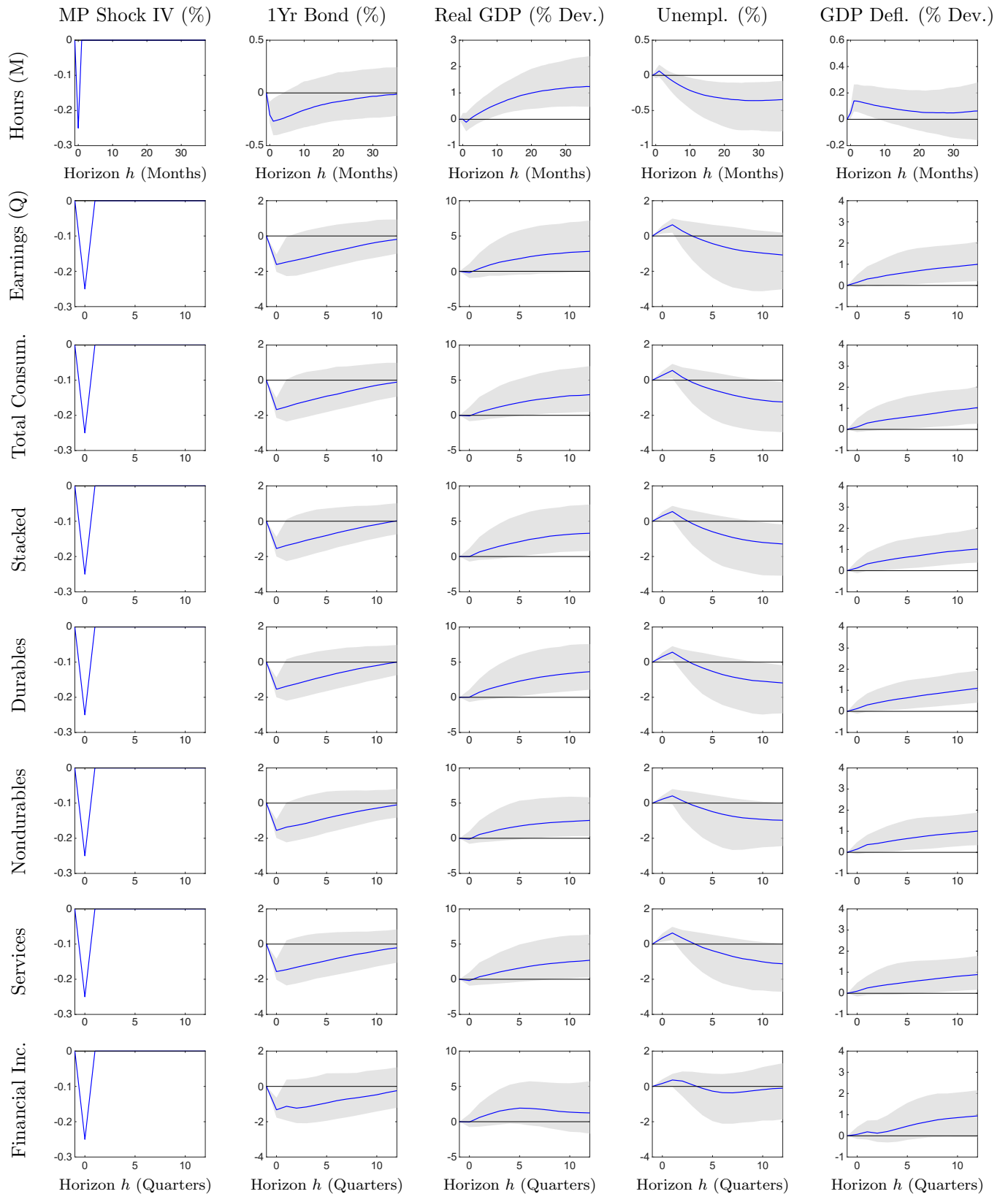
Qualitatively, the aggregate responses from the quarterly models are similar to the responses from the monthly models. Quantitatively, there is a difference in magnitude. The reason is that a 25 basis point surprise at monthly frequency is only roughly a third of a 25 basis point surprise over an entire quarter. In calculations not reported in the paper we time aggregated the responses from the monthly earnings VAR to quarterly frequency and re-scaled the IRFs such that the instrument moves by 25 basis points over the quarter. After this re-scaling, the magnitude of the aggregate responses from the VARs estimated with monthly and quarterly data, respectively, were indeed very similar.

Figure A-1: Responses of Aggregate Variables to Monetary Policy Shock



Notes: Responses to a 25 basis point monetary policy shock. The Aggregate VAR sets $W_t = Y_t$, whereas the functional VAR uses $W_t = [Y_t', \hat{\alpha}_t']'$. The system is in steady state at $h = -1$ and the shock occurs at $h = 0$. The plots depict 10th (dashed), 50th (solid), and 90th (dashed) percentiles of the posterior. GDP defl. and real GDP responses are percentage deviations from the steady state, whereas the other responses are absolute percentages.

Figure A-2: Responses of Aggregate Variables to Monetary Policy Shock

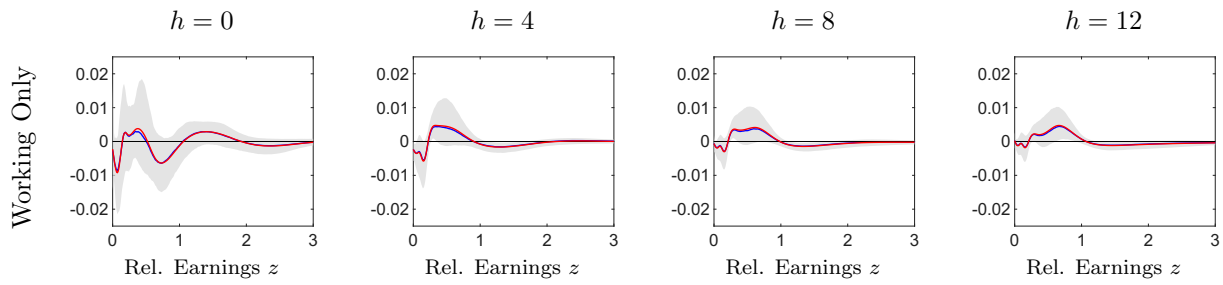
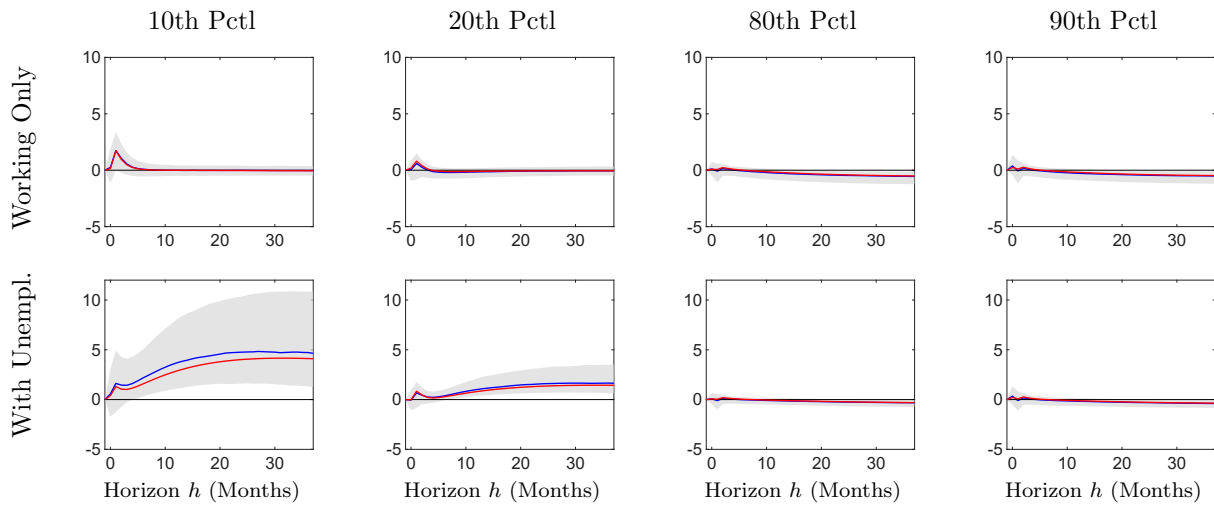


F.2 Responses at the Posterior Mean Parameters

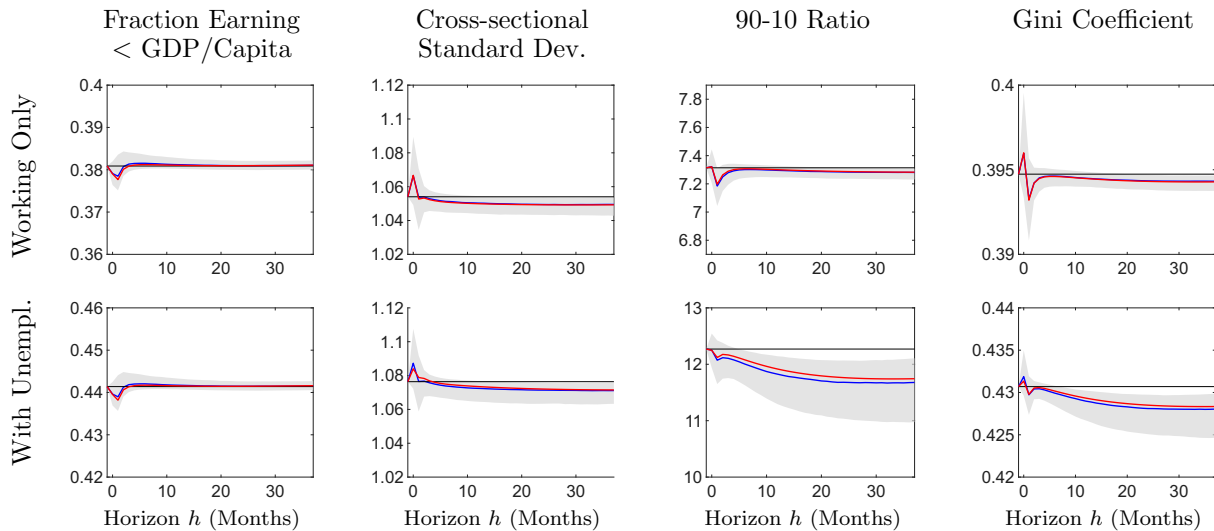
In Footnote 14 we mentioned that the pointwise posterior medians of the density differentials do not generate the posterior medians of the percentile responses. Figure A-3 overlays the responses generated from the posterior mean parameter estimates, which look very similar to the posterior medians reported in Figure 5 of the paper.

Figure A-3: Response of Labor Earnings Distribution to Monetary Policy Shock

Panel (i): Density Differentials

Panel (ii): Rel. Earnings z Percentiles (% Dev. from Steady State)

Panel (iii): Inequality Measures

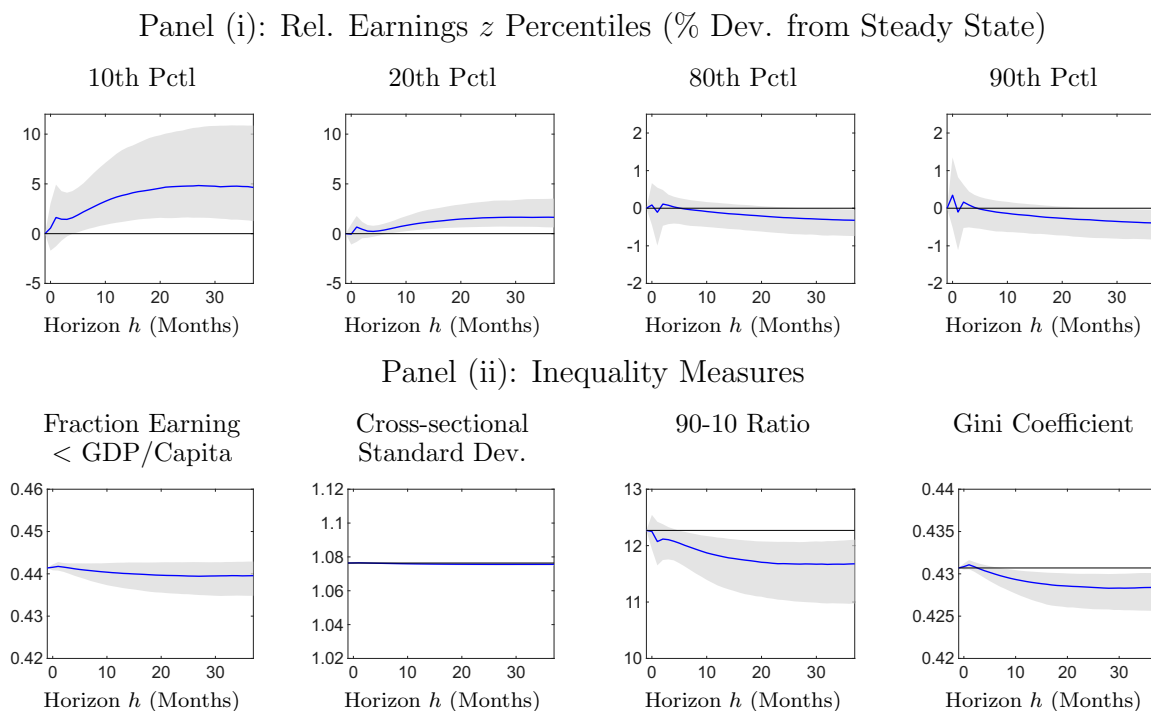


Notes: 10th (dashed), 50th (solid), and 90th (dashed) percentiles of the posterior distribution. The red line corresponds to IRFs that are generated by fixing the parameters at their posterior mean estimates.

F.3 Effect of Unemployment on Earnings Percentile and Inequality IRFs

To justify the claim that the earnings percentile and inequality IRFs are dominated by changes in unemployment we conduct the following exercise. When computing the response of percentiles and inequality statistics we keep the shape of the continuous part of the earnings distribution fixed at its initial level. We only change its normalization to $1 - UR_h$, where UR_h is the unemployment rate response at horizon h . We then compute the percentiles and inequality statistics from the re-normalized continuous part plus the UR_h point mass at zero. The results are plotted in Figure A-4. The percentile responses are visually identical to the “With Unempl.” plots in Figure 5 of the main text, which also account for movements in the shape of the continuous part of the earnings distribution. Accordingly the response of the 90-10 ratio is identical as well. The main feature that the simplified calculation does not capture is the slight reduction in the cross-sectional standard deviation.

Figure A-4: Earnings Percentile and Inequality Responses, Holding Continuous Part Fixed



Notes: The system is in steady state at $h = -1$ and a 25bp monetary shock occurs at $h = 0$. 10th (dashed), 50th (solid), and 90th (dashed) percentiles of the posterior distribution. Horizon h is on the x -axis.

F.4 IRFs from a Model that Stacks Earnings and Consumption Densities

F.4.1 Baseline Analysis

In this section we are comparing the density responses for earnings and consumption from a functional VAR that only includes a single cross-sectional density to a functional VAR that includes two cross-sectional densities simultaneously. The responses are quantitatively very similar. They are plotted in Figures A-5 and A-6.

Figure A-5: Response of Labor Earnings Distribution to Monetary Policy Shock:
Stacked vs. Single Density

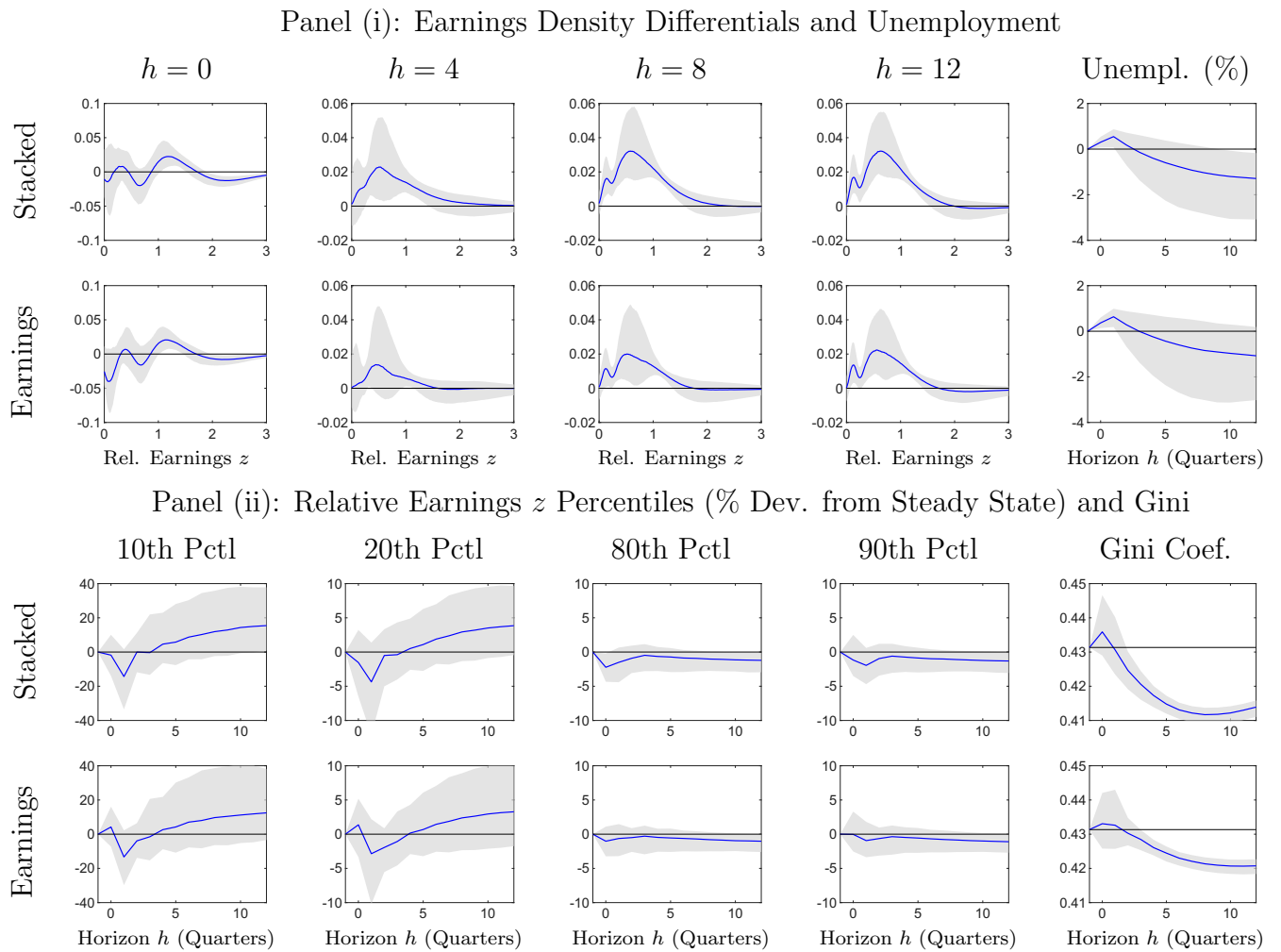
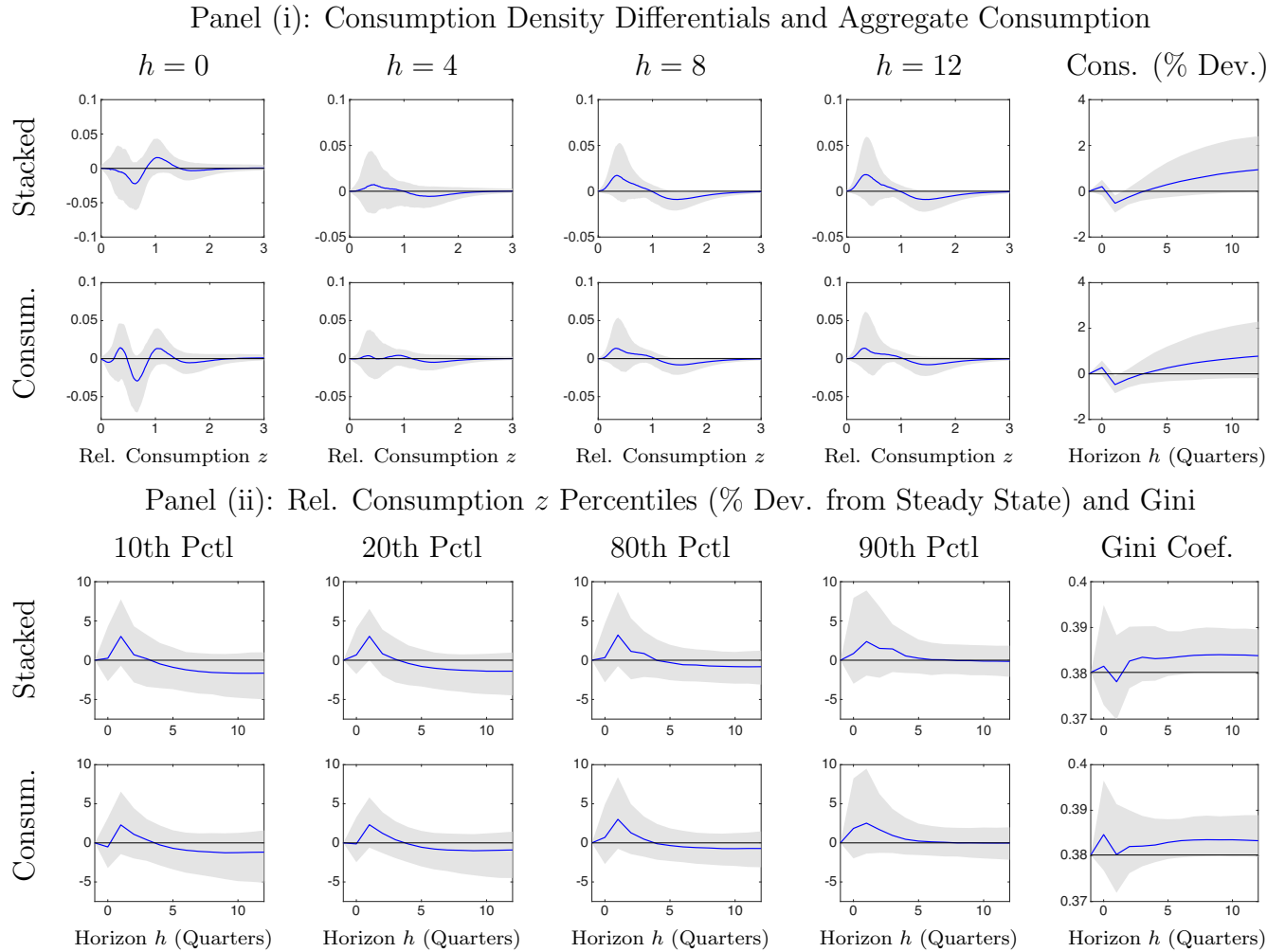


Figure A-6: Response of Consumption Distribution to Monetary Policy Shock:
Stacked vs. Single Density



References

- CHANG, M., X. CHEN, AND F. SCHORFHEIDE (2024): “Heterogeneity and Aggregate Fluctuations,” *Journal of Political Economy*, 132(12), 4021–4067.
- FEDERAL RESERVE BANK OF ST. LOUIS (2022): “Federal Reserve Economic Data (FRED),” St. Louis, MO: Federal Reserve Bank of St. Louis, <https://fred.stlouisfed.org>.
- JAROCINSKI, M., AND P. KARADI (2020a): “Data for ”Deconstructing Monetary Policy Surprises - The Role of Information Shock”,” <https://www.openicpsr.org/openicpsr/project/231538/version/V1/view>, American Economic Association; distributed by Inter-university Consortium for Political and Social Research, Ann Arbor, MI.

——— (2020b): “Deconstructing Monetary Policy Surprises - The Role of Information Shocks,” *Amercian Economic Journal: Macroeconomics*, 12, 1–43.

NATIONAL BUREAU OF ECONOMIC RESEARCH (????): “NBER Based Recession Indicators for the United States from the Peak through the Trough [USRECM],” Retrieved from FRED, Federal Reserve Bank of St. Louis.

PIKETTY, T., E. SAEZ, AND G. ZUCMAN (2018a): “Data for ”Distributional National Accounts: Methods and Estimates for the United States”,” <https://gabriel-zucman.eu/usdina/>, Oxford University Press; distributed by Gabriel Zucman.

——— (2018b): “Distributional National Accounts: Methods and Estimates for the United States,” *Quarterly Journal of Economics*, 133(2), 553–609.

STOCK, J. H., AND M. W. WATSON (2010): “Monthly GDP and GNI – Research Memorandum,” *Manuscript, Princeton University*.

U.S. BUREAU OF ECONOMIC ANALYSIS (2021): “Personal Consumption Expenditures [PCE],” Retrieved from FRED, Federal Reserve Bank of St. Louis.

U.S. BUREAU OF LABOR STATISTICS (2017): “Consumer Expenditure Survey, Public Use Micro Data, 1990-2016,” Washington, DC: United States Department of Labor, https://www.bls.gov/cex/pumd_data.htm#stata.

——— (2021): “Civilian Noninstitutional Population [CNP16OV],” Retrieved from FRED, Federal Reserve Bank of St. Louis.

U.S. CENSUS BUREAU, AND U.S. BUREAU OF LABOR STATISTICS (2022): “Current Population Survey, Basic Monthly, 1990-2016,” Washington, DC: U.S. Department of Commerce and U.S. Department of Labor, <https://www.nber.org/research/data/current-population-survey-cps-basic-monthly-data>, Public-use microdata files.