

Estimated Monthly National Accounts for the United States

Philip Barrett

WP/25/134

IMF Working Papers describe research in progress by the author(s) and are published to elicit comments and to encourage debate.

The views expressed in IMF Working Papers are those of the author(s) and do not necessarily represent the views of the IMF, its Executive Board, or IMF management.

**2025
JUL**



IMF Working Paper

Western Hemisphere Department

Estimated Monthly National Accounts for the United States**Prepared by Philip Barrett***

Authorized for distribution by Nigel Chalk

July 2025

IMF Working Papers describe research in progress by the author(s) and are published to elicit comments and to encourage debate. The views expressed in IMF Working Papers are those of the author(s) and do not necessarily represent the views of the IMF, its Executive Board, or IMF management.

ABSTRACT: I jointly estimate monthly series for GDP and eight subcomponents for the US since 1950. The series match 1) quarterly national accounts equivalents, 2) exact data on monthly consumption, and 3) past relationships with other monthly indicators. I estimate the Kalman filter parameters by GMM, allowing fast calculation of confidence intervals for monthly estimates including parameter uncertainty, and validate the confidence intervals. After 1970 standard errors are tight, less than 0.3pp of GDP, and point estimates informative, with standard deviations four times the standard error. I provide confidence intervals for recessions and show that output peaks line up well with the onset of NBER recessions, but troughs often predate NBER equivalents.

JEL Classification Numbers:	E01, C32
Keywords:	Kalman Filter; GDP; recession; GMM
Author's E-Mail Address:	pbarrett@imf.org

1 Introduction

Many important macroeconomic data series for the United States are available at monthly frequency. Labor market data, including unemployment, vacancies, hours worked, and some wage series, are all monthly, as are detailed data on consumer and producer prices. So too are data on trade, financial markets (including interest rates, exchange rates, and the like), federal government accounts, housing starts, industrial production, and many more. Yet perhaps the single most important summary measures of economic performance—GDP and its expenditure components—are only available quarterly.

This gap in the data forces researchers to make a particularly frustrating choice. One option is to use monthly data in their regressions or models, but then their results are not directly comparable to the national accounts. Alternatively, one can use quarterly data, aggregating where required. This way, GDP is included, but at the cost of throwing away much of the other data, including potentially identifying variation.

This situation is made all the more frustrating by the fact that many monthly data series are very close, and in some cases exact, substitutes for many GDP components. This is clearest for consumption, where the quarterly national accounts series is just the three-month sum of the monthly personal consumption expenditure. For imports, exports, and inventories monthly data are close analogues of the quarterly national accounts series, but the correspondence is not exact and coverage is partial. Yet for other the components, i.e. fixed investment and government spending, there are no such equivalent data.

The preceding suggests that it might not only be very useful to have a set of monthly series directly comparable to GDP and its major constituents, but also that it should be feasible to so.

This paper makes three contributions to tackling this problem. The first is to provide a set of monthly point estimates for GDP and nine subcomponents¹ from January 1950 to March 2025. I derive these estimates using a Kalman filter, which allows me to include information using three types of data: the quarterly national income and product accounts accounts (NIPA); exact monthly data for some compo-

¹These are: private goods consumption, private services consumption, residential fixed investment, non-residential fixed investment, inventories, imports, exports, and government spending.

nents; and other monthly indicators. This latter category is, in principle, a broad one and could include anything from near-exact monthly estimates of (e.g. monthly trade data from the Census Bureau) to more distantly-related ones (e.g. housing starts). The resulting estimates add up to the quarterly NIPA data almost exactly.

A key novelty of this paper is that I use a generalized method of moments (GMM) approach to estimate the parameters of the Kalman filter. Usually, Kalman filters are estimated by maximum likelihood, but in this application the parameter space grows quickly enough that this would suffer from an acute curse of dimensionality.² Instead, I exploit a particular feature of the data in this setting: that the NIPA data provide quarterly averages of all the hidden states. This means that one can compute *any* moment of the quarterly average of the data exactly and choose parameters to set the equivalent theoretical moments equal to the values of the data.

The resulting moment conditions can be divided into two subsets: those defining parameters governing the dynamics of the GDP components, and everything else. The first group is over-identified by a subset of moments of the data. Conditional on estimated of the first group of parameters, all the other parameters are just-identified. This makes GMM fast and easy: one can estimate the parameters governing the dynamics of each GDP component one-by-one, then compute the remaining parameters as the exact values which solve their relevant moment conditions.³

There is a cost to using GMM: moving average components at higher-than-quarterly frequency cannot be identified from this data. The stance in this paper is that this is price worth paying for the speed and simplicity of GMM.

The second major contribution of this paper is to provide confidence intervals for the monthly point estimates. These are narrow. The estimated standard error for monthly GDP is less than 0.3 percent from 1970, and falls further in later decades. At the same time, the monthly point estimate captures considerable month-to-month variation in economic activity relative to its precision. From 1970 onward, the one-

²The baseline model has 60 parameters, and versions used in the various robustness checks push this higher, making maximum likelihood difficult. This also suggests an answer to an obvious question: why hasn't someone else done this already? One possibility is that ML estimation is just too painful. Indeed, Aruoba et al. (2009) note the computational expense of estimating the Kalman likelihood in a similar setting.

³This is only true for unweighted GMM. For efficient GMM, the moments are not separable in this way and one must minimize the weighted moment errors together. In practice, this does not matter. The sequentially-estimated parameters provide an excellent starting guess for efficient GMM, and the gains from efficiency are very small.

period-ahead standard deviation in the GDP point estimate is more than four times its standard error. Post-2000, this ratio rises to around seven.⁴

The confidence intervals are themselves the sum of two parts, reflecting two sources of uncertainty. The *filter uncertainty* arises because the signals are noisy, and remains even if the model parameters are known perfectly. However, the model parameters are not known; they are estimated. And so the estimated monthly GDP components are also subject to *estimation uncertainty*. Filter uncertainty is straightforward to calculate, and is computed using the error covariance of the Kalman smoother. Estimation uncertainty is given by the covariance of the GDP component point estimates, which can be computed using the delta method. Adding the two source of uncertainty gives time-varying confidence intervals for each GDP components.

The confidence intervals also illuminate the drivers of uncertainty over monthly GDP. Because it reflects the information content of the noisy signals given the parameters, filter uncertainty is typically constant over periods of time, dropping in steps as new data become available. This is most obvious when exact consumption data become available in January 1959. Of course, the filter uncertainty over the consumption series drops to zero. But because I allow for correlated innovations in the GDP components, the filter uncertainty for other series also falls, by around half. For example, residential investment and goods consumption are positively correlated at quarterly frequency, since when people build new houses they also tend to buy refrigerators, televisions, and such. By exploiting this correlation, the Kalman filter can thus use the increased precision of the consumption series to improve the precision of the residential investment estimate. Estimation uncertainty, in contrast, is time-varying and increases when monthly indicators are more variable, since the uncertainty over the model coefficients is multiplied by the variation in the data.

I also apply this measurement of uncertainty to create confidence intervals for peaks and troughs in output. I compare these to NBER recession dates and find that output peaks tend to align more closely with the NBER start dates than the troughs do. Moreover, differences in dating peaks largely reflect “false start” recessions, where output shows a slight dips before a more decisive decline a few quarters later. In such cases, disagreements between different measures are not material. However, the disagreements on the end dates of the recession are more systematic and more often statistically significant, with the monthly GDP measure tending to call an end to

⁴The signal-to-noise ratio is usually defined as the square of the numbers cited here.

recessions earlier than the NBER dates. I also decompose business cycle variation in monthly GDP, finding that quarterly NIPA data understate the variation in monthly GDP due to government spending prior to the 1990s.

The third major contribution is to provide evidence that the confidence intervals described above are valid. The ideal test would be to check the coverage ratios of my confidence intervals against the data. This is impossible since the monthly data are unobserved. As a second best, I construct an alternate version of my estimates where I re-run the entire process pretending that one of the known monthly consumption series is unobserved. I re-estimate all the parameters using only some noisy monthly indicators instead (real retail sales or durable goods consumption). This aims to mimic what I do for the other GDP components, where there is no exact data and only noisy indicators. But because the consumption data do exist, I can assess the confidence intervals, checking if $x\%$ of outcomes fall inside any given $x\%$ confidence interval. The confidence intervals are almost exactly correct for any x . Although proving it is impossible, this is at least strong suggestive evidence that this method can correctly characterize the distribution of the unknown monthly series. This exercise also highlights the importance of including estimation uncertainty. Without filter uncertainty alone, the confidence intervals for the goods consumption series are much too narrow.

Related Literature. There is a rich literature producing monthly estimates of real activity for the United States. The earliest versions (Stock and Watson (1988), Stock and Watson (1989)) were factor models which aimed to produce activity indices from common components in various monthly data series. Mariano and Murasawa (2003) were one of the first to produce an index comparable to real GDP data by using a state space model to impose quarterly adding-up constraints, something almost universally imposed in subsequent works (including this one). Aruoba et al. (2009) use a similar framework, estimating a Kalman Filter by maximum likelihood to produce estimates of daily GDP, using high-frequency observable data. These papers typically use rather small sets of observable data, most commonly only four series plus quarterly GDP.

Stock and Watson (2010) and especially Brave et al. (2019) use much larger sets of observable data in state space frameworks to extract more potential information about monthly GDP, with the latter using over 500 different monthly variables. Stock and Watson (2010) also outline methods for estimating both GDP and GNI separately,

an idea developed further in Koop et al. (2023) who explicitly reconcile expenditure and output measures of GDP in a single framework. In contrast, I focus on providing monthly equivalents to the expenditure-side accounts only.

There is also a large literature seeking to increase the frequency of data in low-income countries using alternate data sources. These include work aiming to produce quarterly real activity measures in countries where only annual data is available (see Stanger (2020), and Akbal et al. (2023) among others), as well as work on extending and validating national accounts using satellite imagery (most notably Hu and Yao (2022) and Beyer et al. (2022)).

Outside of the academic literature, Standard and Poor’s also estimate monthly GDP, with the headline series freely available. However, little information is available on exactly how this is computed beyond a brief heuristic description.

This paper relates the statistical model to the data in Section 2, introduces the data in Section 3, describes the GMM estimation strategy in Section 4, presents the main results in Section 5, and validates those results in Section 6. Section 7 concludes.

2 Set up

2.1 Overview

I assume that real monthly GDP can be written as the sum of N components:

$$gdp_t = \sum_{i=1}^N y_t^i \tag{1}$$

where t ranges from 1 to T . In the baseline specification, $N = 8$ and the components are goods consumption, services consumption, residential investment, non-residential investment, imports, exports, government consumption & investment, and changes in inventories.⁵

⁵Strictly speaking, this does not sum exactly the GDP, since chain-linking means that there is an aggregation in the quarterly NIPA data. This is close to zero near to the base year, 2017, but is larger earlier in the sample. To avoid extensive modeling of the aggregation error, my baseline results abstract from this, omitting the errors and defining “GDP” as the sum of the components. In robustness checks and in the data made available with this paper, I also include outputs including the aggregation error.

Letting $y_t = (y_t^1, \dots, y_t^N)'$, I aim to estimate a vector of conditional means and variances:

$$\mu_t = \mathbb{E}(y_t|I_T) \quad \Sigma_t = \text{Var}(y_t|I_T)$$

The conditioning set I_T consists of observations $t = 1, \dots, T$ for three sources of information. Also summarized in Table 1, these sources are:

1. *Quarterly NIPA data.* We denote these Y_t^i for $i = 1, \dots, N$ and write $Y_t = (Y_t^1, \dots, Y_t^N)'$. For $t = 3, 6, 9, \dots$, we observe:

$$Y_t^i = \sum_{l=0}^2 y_{t-l}^i \quad (2)$$

2. *Exact monthly data for some components.* For components $i = 1, \dots, K$ we have exact monthly data, at least in some periods. The most prominent such series is consumption, which comes from the monthly BEA income and outlays report (see Section 3 for other cases). These “hard” data are denoted x_t^i , with $x_t = (x_t^1, \dots, x_t^K)'$ the vector of these observations.
3. *Other monthly indicators correlated with y_t^i .* For example, monthly federal spending, government employment, partial inventories or investment data, and such. We denote the time- t vector of such indicators for component i as q_t^i , the full set of indicators as $q_t = (q_t^1, \dots, q_t^{N'})'$, and the total number of indicators as P .

Component index	Quarterly NIPA data	Exact monthly data	Monthly indicators
$i = 1, \dots, K$	✓	✓	Possibly
$i = K + 1, \dots, N$	✓	✗	✓

Table 1: Data structure

Given estimates for the conditional mean and variances of the component, the conditional mean and variance for monthly GDP is given by:

$$\mathbb{E}(gdp_t|I_T) = 1'_N \hat{\mu}_t \quad \text{Var}(gdp_t|I_T) = 1'_N \hat{\Sigma}_t 1_N \quad (3)$$

2.2 The Kalman form

I use a Kalman filter to estimate μ_t, Σ_t . The Kalman Filter assumes that a vector of hidden states ξ_t and observables Z_t are generated by stochastic process:

$$\begin{aligned}\xi_t &= F_t \xi_{t-1} + V_t & \text{cov}(V_t) &= Q_t \\ Z_t &= H_t \xi_t + W_t & \text{cov}(W_t) &= R_t\end{aligned}$$

The key assumption underpinning the Kalman filter is that the state and observation errors V_t and W_t are each serially uncorrelated.

The rest of this section is devoted to defining 1) the mapping of the elements of this problem into ξ_t and Z_t , and 2) the structure of the matrices F_t, H_t, R_t and Q_t .

2.3 The state equations

For each i , I assume that the data generating process for y_t^i is given by departures from a long-run polynomial growth path which can be modeled as an AR(L)⁶ process with correlated errors:

$$y_t^i = \bar{y}_t^i (1 + \hat{y}_t^i) \tag{4}$$

$$\log \bar{y}_t^i = \sum_{m=0}^M \gamma_m^i t^m \tag{5}$$

$$\hat{y}_t^i = \sum_{l=1}^L \rho_l^i \hat{y}_{t-l}^i + w_t^i \tag{6}$$

The assumption that the trend growth path is polynomial is merely a flexible way to capture (possibly differing) trends in the GDP components. Since this is an exercise of interpolation, not forecasting, we can sidestep questions over stationarity and cointegration. Here, the main role of the trend in this case is just to provide a rescaling which means that the stochastic part of y_t^i is plausibly covariance-stationary.

To first order, defining $\hat{y}_t^i = \left(\frac{\bar{y}_t^i}{\bar{y}_t^i} - 1 \right)$ is the same as defining it as $\hat{y}_t^i = \log \left(\frac{\bar{y}_t^i}{\bar{y}_t^i} \right)$. However, using ratios is preferable because 1) it makes imposition of the quarterly

⁶In the exposition, I assume that all monthly GDP components series have the same lag length. However, when estimating the model I relax this assumptions. Extending the notation to reflect this is straightforward but ugly, so is omitted here.

adding-up constraints a little easier, and 2) changes in inventories can be negative.⁷

Although the Kalman form requires that the w_t^i are serially uncorrelated, I allow for cross-correlation across the monthly GDP components. That is, $Cov(w_t^i, w_t^j) \equiv \sigma_{i,j}^2$ may be different from zero. Such cross-correlations are potentially an important source of information. For example, one might expect abnormally high consumption to be correlated with both increased imports and reduced inventories.

In terms of the Kalman form above, this means that ξ_t is merely the L stacked lags of $\hat{y}_t$, F_t is the appropriate matrix of ρ^i coefficients, and Q_t is an $LN \times LN$ matrix which is all zero except for the top-left $L \times L$ block which is the covariance matrix of the w_t^i .

Some more notation is useful when writing down the moment conditions used to estimate the model later. I assume $\hat{y}_t^i$ is invertible with moving average representation:

$$\hat{y}_t^i = \sum_{l=0}^{\infty} a_l w_{t-l}^i \quad (7)$$

And the covariances of $\hat{y}_t^i$ are given by:

$$Cov(\hat{y}_t^i, \hat{y}_{t-l}^i) = \theta_l^i \quad l = 0, \dots, \infty \quad (8)$$

2.4 The observation equations

Given that there are three distinct sources of information, Y_t, x_t, q_t , it will be helpful to divide the elements of the Kalman observation equations Z_t , H_t , and R_t into three parts:

$$\begin{aligned} Z_t &= \begin{bmatrix} Z_t^Y \\ Z_t^x \\ Z_t^q \end{bmatrix} \\ &\begin{matrix} N \times 1 \\ K \times 1 \\ P \times 1 \end{matrix} \\ &\begin{matrix} (N+K+P) \times 1 \end{matrix} \end{aligned} \quad \begin{aligned} H_t &= \begin{bmatrix} H_t^Y \\ H_t^x \\ H_t^q \end{bmatrix} \\ &\begin{matrix} N \times LN \\ K \times LN \\ P \times LN \end{matrix} \\ &\begin{matrix} (N+K+P) \times LN \end{matrix} \end{aligned} \quad \begin{aligned} R_t &= \begin{bmatrix} R_t^Y & 0 & 0 \\ 0 & R_t^x & 0 \\ 0 & 0 & R_t^q \end{bmatrix} \\ &\begin{matrix} N \times N \\ K \times K \\ P \times P \end{matrix} \\ &\begin{matrix} (N+K+P) \times (N+K+P) \end{matrix} \end{aligned}$$

In this section, I define the structure of each of these submatrices.

⁷Similarly, I also assume that the trend for inventories is polynomial in the level, not the log.

2.4.1 Quarterly NIPA data

The quarterly NIPA aggregation constraint can be written as:

$$Y_t = \sum_{l=0}^2 \left(\frac{\bar{y}_{t-l}^i}{\bar{Y}_t^i} \right) (1 + \hat{y}_{t-l}^i)$$

$$\Rightarrow \hat{Y}_t^i = \sum_{l=0}^2 \left(\frac{\bar{y}_{t-l}^i}{\bar{Y}_t^i} \right) \hat{y}_{t-l}^i$$

where $\hat{Y}_t = \frac{Y_t}{\bar{Y}_t} - 1$ is the deviation of quarterly NIPA data from trend and $\bar{Y}_t^i = \sum_{l=0}^2 \bar{y}_{t-l}^i$. Thus, $Z_t^Y = \hat{Y}_t$ with missing observations for two months in every quarter (the Kalman filter can handle missing data very easily). The state loading matrix H_t^Y has entries given by the appropriate $\frac{\bar{y}_{t-l}^i}{\bar{Y}_t^i}$. And since the aggregation equations hold exactly, R_t^Y is zero (although in practice, numerical stability requires using $R_t^Y = \eta I_{N \times N}$ where η is a very small number).

2.4.2 Exact data

Likewise, the observations for the exact data can be written as:

$$\frac{x_t^i}{\bar{y}_t^i} - 1 = \hat{y}_t^i$$

Thus $Z_t^x = (x_t^1/\bar{y}_t^1 - 1, \dots, x_t^N/\bar{y}_t^N - 1)'$, H_t^x has ones along the primary diagonal and zeroes elsewhere, and R_t^x is also zero.

2.4.3 Other monthly indicators

I denote by $q_t^{i,j}$ the j -th element of q_t^i , and fit each $q_t^{i,j}$ to its own polynomial time trend, $\bar{q}_t^{i,j}$ where:

$$\log \bar{q}_t^{i,j} = \sum_{m=0}^M \eta_m^{i,j} t^m \quad (9)$$

And departures from trend are: $\hat{q}_t^{i,j} = \frac{q_t^{i,j}}{\bar{q}_t^{i,j}} - 1$.

To eliminate serial correlation, I estimate ARIMA processes for each monthly indicator $\hat{q}_t^{i,j}$ with lag length chosen by AIC, and denote the serially uncorrelated

residuals by $\epsilon_t^{i,j}$. I assume that these are related to the innovations in $\hat{y}_t^i$ by:

$$\epsilon_t^{i,j} = \kappa^{i,j} w_t^i + u_t^{i,j} \quad (10)$$

where $u_t^{i,j}$ is white noise and $cov(w_t^i, u_t^{i,j}) = 0$. Concerns about the validity of this last assumption are not relevant here, since equation (10) expresses a statistical relationship and not a causal one. It just aims to capture the systematic correlation of the data with the state. Indeed, the orthogonality of w_t^i and u_t^i defines the decomposition of $\epsilon_t^{i,j}$ into a part correlated with w_t^i and a part which is a residual.

Since the left hand side of equation (10) is just data, it can easily be written as:

$$\epsilon_t^{i,j} = \kappa^{i,j} \hat{y}_t^i - \sum_{l=1}^L \kappa^{i,j} \rho_l^i \hat{y}_t^i + u_t^{i,j} \quad (11)$$

Then Z_t^q is the stacked vector of the $\epsilon_t^{i,j}$ and H_t^q is the appropriate matrix of the $\kappa^{i,j}$ or $\kappa^{i,j} \rho_l^i$. Both the $\kappa^{i,j}$ terms and the full set of entries of R_t^q are unknown and will be estimated by GMM as described in Section 4. To avoid the size of the parameter space growing with the square of the number of indicators, I let R_t^x be block-diagonal, restricting the cross-correlation of monthly indicators for different components to zero.

An inferior specification of the observation equations. The set-up in equation (10) puts the indicator residual on the left hand side and hidden state on the right hand side. This might seem like it is the wrong way round. The reader may be wondering: shouldn't the NIPA component be on the left, and the indicator on the right? That argument is all the more compelling if the indicator is itself a component of the hidden state: for example, if $\hat{y}_t^i$ were the export component of GDP and $\epsilon_t^{i,j}$ were the residual on monthly goods exports. And with $\hat{y}_t^i$ on the left, one could put *all* the relevant indicators on the right hand side, like a multiple regression. This would also save on a lot of extra equations. Such a set-up would look something like this:

$$w_t^i = \sum_j \beta^{i,j} \epsilon_t^{i,j} + \tilde{u}_t^{i,j} \quad (12)$$

This is a feasible alternative to equation (11), but it has no clear advantage over it. Since equation (11) is a purely statistical relationship rather than a causal one, we cannot give a different interpretation to one side over the other. This is just

two variables covarying, not one driving the other. And by allowing for non-zero cross-equation covariances in R_t^x , the Kalman filter's inference about y_t^i is a function of multiple monthly indicators, substituting for the direct role played by the linear combination of indicators on the right hand side of equation (12). Indeed, the same covariation in the $\epsilon_t^{i,j}$ determines the cross-equation residual covariance from equation (10) as would show up in the denominator of $\beta^{i,j}$ when expressed as a regression coefficient.

However, specifying the observation equations as in equation (10) rather than (12) has a key advantage: it handles missing observations much better. If different monthly indicators are only available for different subsamples of the data, the Kalman filter handles this by dropping the appropriate restrictions from the filter algorithm in those periods. With a set-up like equation (10) this permits dropping just the missing subset of monthly indicators. But if the observation equations are specified as in equation (12), this would involve dropping all observations where *any* monthly indicator is missing.

An incorrect specification of the observation equations. An important property of equation (10) is that it is consistent with zero serial correlation of the error terms. Other specifications can easily violate this. For example, consider the alternative observation equation:

$$\hat{q}_t^{i,j} = \alpha^{i,j} + \beta^{i,j} \hat{y}_t^i + \tilde{u}_t^i \quad (13)$$

This is somewhat simpler than equation (10), and although equation (13) has the advantage that it is easy to estimate⁸ it comes at the cost of producing entirely incorrect estimates of $\hat{y}_t^i$. This occurs because the residuals on this observation equation will also be serially correlated, in violation of the assumptions of the Kalman filter⁹. Thus, using equation (13) as one of the observation equations would mis-attribute the timing of the information in the monthly indicator $\hat{q}_t^{i,j}$, assigning all variation to $\hat{y}_t^i$, when instead this should be apportioned across $\hat{y}_t^i$ and its lags.

⁸In that a regression of $\frac{1}{3} \sum_{l=0}^2 \hat{q}_{t-l}^{i,j}$ on Y_t^i would provide consistent estimates of $\alpha^{i,j}, \beta^{i,j}$.

⁹Except in the knife-edge case where the lag structure of $\hat{q}_t^{i,j}$ is *exactly* the same as that of $\hat{y}_t^i$.

3 Data

3.1 Quarterly National Income and Product Accounts

Quarterly data come from the national income and product accounts (NIPA). These are available back to 1950Q1. This defines the starting data for my monthly GDP series, January 1950. These data are shown in Figure 1, which also includes the estimated trends (on which, more later).

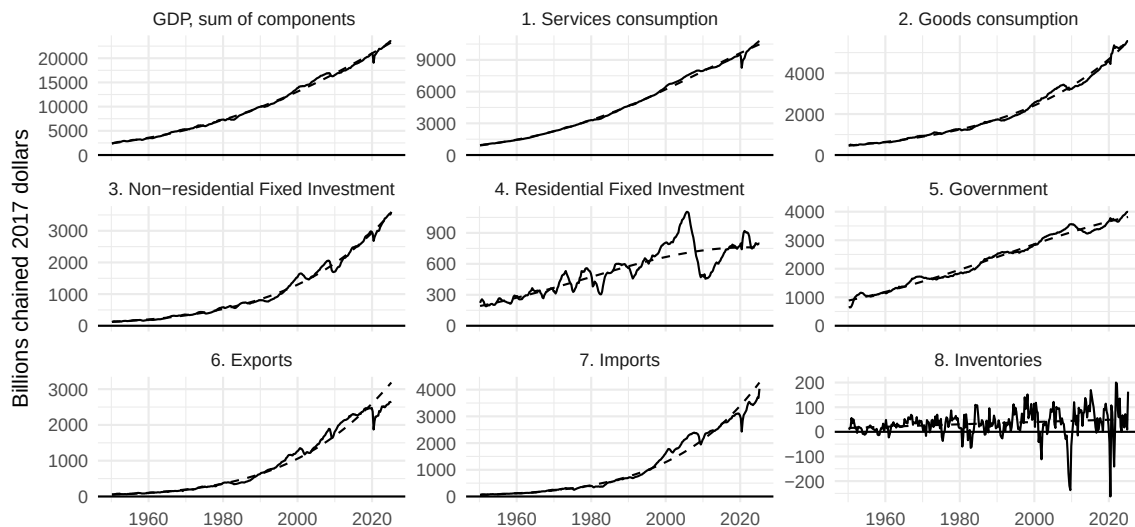


Figure 1: Quarterly NIPA data

Figure shows quarterly National Income and Products Account Data (solid line) with estimated trends (dashed lines). Sample: 1050Q1-2025Q1

3.2 Monthly indicators

The BEA publishes monthly data for two of the GDP components: consumption of goods and services. These are the direct monthly analogues of the corresponding quarterly NIPA series, and add up to them. I use eight further monthly data series as indicators for the remaining GDP components. These are summarized in Table 2 and their time series shown in Figure 2. The residualized versions are shown in Figure A1 in Appendix A.

The selection of the monthly indicators is an important choice in constructing a monthly GDP series. A full discussion of model selection is postponed to section

6.3. However, as a general principle, there should be a high bar to including extra series as indicators, since adding extra data series will often *worsen* the accuracy of the estimated monthly GDP series. Adding extra data series reduces filter uncertainty but increases estimation uncertainty.¹⁰ When the relationship between the data and the hidden state is poorly estimated, the increased estimation uncertainty will dominate. As we do not actually see the hidden state, this is likely the case. Just throwing extra data series at the model may make estimated monthly GDP worse.

An important exception is cases where there are strong a priori reasons to believe that the monthly indicators are very close substitutes for the GDP components or subset of them. For example, the Census Bureau’s monthly import and export series are not only conceptually very similar to the BEA’s quarterly equivalents – with slight differences in coverage – but are in fact inputs to the BEA’s calculations. Likewise, the BEA’s own monthly change in inventories series is very close to being exact data for monthly inventories. If this were available prior to 1997, it would be the only data series necessary for estimating monthly inventories. Since it is not, I also include the change in manufacturing inventories, which covers only part of total inventories. For the investment and government spending series, there are no similarly obvious monthly analogues.

y_t^i	x_t^i	$q_t^{i,j}$	Starting
Goods Consumption	✓		Jan 1959
Services Consumption	✓		Jan 1959
Non-residential fixed investment	✗	Non-residential construction	Jan 2002
Residential fixed investment	✗	Housing starts	Jan 1960
		Single-family housing starts	Jan 1959
Government spending	✗	None	
Exports	✗	Census export volume	Jan 1992
Imports	✗	Census import volume	Jan 1992
Inventories	✗	Change in manufacturing inventories	Jan 1950
		Change in inventories	Jan 1997
Aggregation error	✗	None	

Table 2: Monthly data series used

¹⁰In Section 5.3, I decompose the uncertainty over the estimated monthly series for GDP and components into contributions from these two sources.

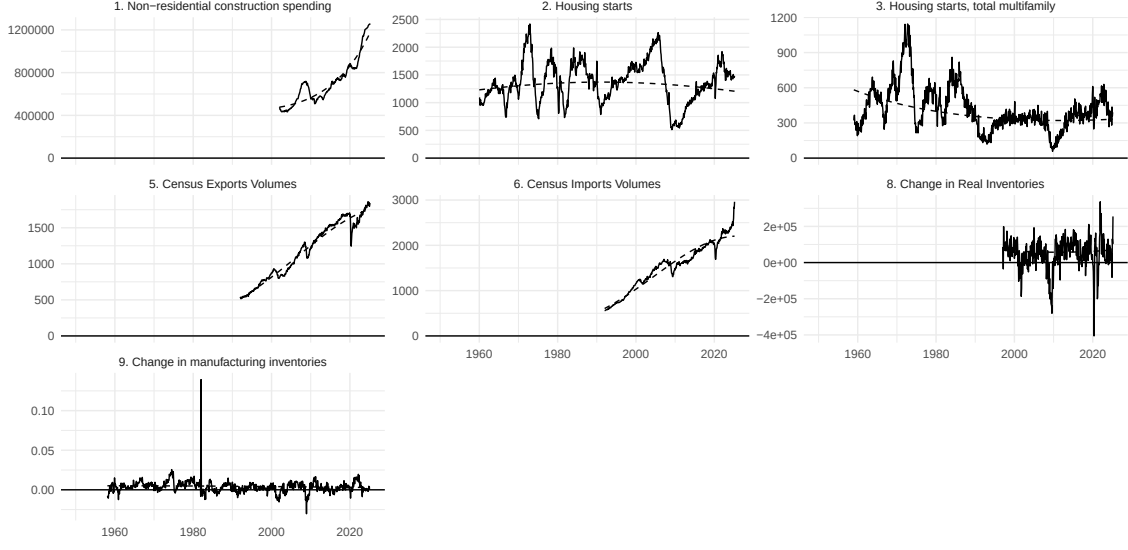


Figure 2: Monthly Indicators

Figure shows monthly indicators (solid line) with estimated trends (dashed lines). Sample: Various-Mar 2025

4 Estimation Streategy

4.1 Trend

To estimate the γ^i time trend parameters, I first regress the quarterly NIPA data Y_t^i on a polynominal time trend:

$$\log Y_t^i = \sum_{m=0}^M \tilde{\gamma}_m^i t^m + e_t^i$$

Then I set:

$$\log \bar{y}_t^i = \frac{1}{3} \sum_{m=0}^M \tilde{\gamma}_m^i (t+1)^m \quad (14)$$

This means that $\bar{Y}_t^i = \bar{y}_{t-1}^i$, and $\bar{Y}_t^i = \frac{1}{3} \sum_{l=0}^2 \bar{y}_{t-l}^i$ holds only to first order. But in practice, the errors are very small. For given M , equation (14) can be inverted easily to recover the γ_m^i in equation (5). In the baseline case, where $M = 2$, this gives:

$$\gamma_0^i = \frac{1}{3} (\tilde{\gamma}_0^i + \tilde{\gamma}_1^i + \tilde{\gamma}_2^i) \quad \gamma_1^i = \frac{1}{3} (\tilde{\gamma}_1^i + 2\tilde{\gamma}_2^i) \quad \gamma_2^i = \frac{1}{3} \tilde{\gamma}_2^i$$

4.2 The state equations

To estimate $\sigma_i, \rho_1^i, \dots, \rho_L^i$ by GMM, I compute the relationship between the autocovariances of $\hat{Y}_t^i$ and those of $\hat{y}_t^i$ using the approximation that $\bar{y}_t^i \simeq \frac{1}{3}\bar{Y}_t^i$:

$$\begin{aligned} Var \hat{Y}_t^i &= \frac{1}{9} Var \left(\sum_{j=0}^2 \hat{y}_{t-j}^i \right) \\ &\simeq \frac{1}{9} (3\theta_0^i + 4\theta_1^i + 2\theta_2^i) \end{aligned} \quad (15)$$

$$\begin{aligned} Cov(\hat{Y}_t^i, \hat{Y}_{t-3k}^i) &= \frac{1}{9} Cov \left(\sum_{j=0}^2 \hat{y}_{t-j}^i, \sum_{j=0}^2 \hat{y}_{t-3k-j}^i \right) \\ &\simeq \frac{1}{9} (\theta_{3k-2}^i + 2\theta_{3k-1}^i + 3\theta_{3k}^i + 2\theta_{3k+1}^i + \theta_{3k+2}^i) \quad k \geq 1 \end{aligned} \quad (16)$$

This is a set of over-identifying moment conditions for $\sigma_i, \rho_1^i, \dots, \rho_L^i$. Since the right-hand side formulas are all functions of $\sigma_i, \rho_1^i, \dots, \rho_L^i$, replacing the autocovariances of $\hat{Y}_t^i$ with their sample analogues gives a set of relationships ready to include in a GMM estimation routine.

Not all stochastic processes for $\hat{y}_t^i$ are identified by this set of moment conditions. For example, an MA(3) process for $\hat{y}_t^i$ is not identified by these moment conditions. Intuitively, matching the moments of quarterly averages will not distinguish between models which differ at frequencies shorter than one quarter. However, strictly autoregressive processes can be identified by their effect on covariances of quarterly averages at sufficiently distant horizons.¹¹

We can also derive a set of moment conditions for the *covariance* of the residuals analytically from the covariances of the $\hat{Y}_t$.

$$Cov(\hat{Y}_t^i, \hat{Y}_t^j) = \frac{1}{9} \left(a_0^i a_0^j + (a_0^i + a_1^i)(a_0^j + a_1^j) + \sum_{s=0}^{\infty} \left(\sum_{l=0}^2 a_{s+l}^i \right) \left(\sum_{l=0}^2 a_{s+l}^j \right) \right) \sigma_{i,j}^2 \quad (17)$$

where $\sigma_{i,j}^2 = Cov(w_t^i, w_t^j)$. Replacing the covariance with its sample analogue gives another moment condition, which we can add to the GMM estimation.

¹¹The identification may not be very good, since long-duration covariances tend all to be close to zero. But that is an empirical matter, and will show up as large standard errors on the parameter estimates.

4.3 Observation equations

To estimate equation (10), I again form moment conditions using the quarterly average data. To start with, I construct the equivalent quarterly residual for the indicator $q_t^{i,j}$ as:

$$E_t^{i,j} = \frac{1}{3} \sum_{l=0}^2 \epsilon_{t-l}^{i,j}$$

By expanding the definitions of $E_t^{i,j}$ and $\hat{Y}_t^i$ we get that:

$$\begin{aligned} Cov(E_t^{i,j}, \hat{Y}_t^i) &= Cov\left(\frac{1}{3} \sum_{l=0}^2 \epsilon_{t-l}^{i,j}, \sum_{l=0}^2 \left(\frac{\bar{y}_t^i}{\bar{Y}_t^i}\right) \hat{y}_{t-l}^i\right) \\ &\simeq \frac{1}{9} (3a_0^i + 2a_1^i + a_2^i) Cov(\epsilon_t^{i,j}, w_t^i) \end{aligned}$$

Then since $Cov(w_t^i, u_t^{i,j}) = 0$, we have that $\kappa^{i,j} = Cov(\epsilon_t^{i,j}, w_t^i)/\sigma_i^2$ Where σ_i^2 is the variance of w_t^i . Substituting in, we get:

$$Cov(E_t^{i,j}, \hat{Y}_t^i) = \frac{1}{9} (3a_0^i + 2a_1^i + a_2^i) \kappa^{i,j} \sigma_i^2 \quad (18)$$

Again, we can replace $Cov(E_t^{i,j}, \hat{Y}_t^i)$ with its sample analogue to get a moment condition suitable for GMM.

To calculate $Cov(u_t^{i,j}, u_t^{k,l})$, we can form another moment condition:

$$\begin{aligned} Cov(\epsilon_t^{i,j}, \epsilon_t^{k,l}) &= Cov(\kappa^{i,j} w_t^i + u_t^{i,j}, \kappa^{k,l} w_t^k + u_t^{k,l}) \\ &= \kappa^{i,j} \kappa^{k,l} \sigma_{i,k}^2 + Cov(u_t^{i,j}, u_t^{k,l}) \\ \Rightarrow Cov(u_t^{i,j}, u_t^{k,l}) &= Cov(\epsilon_t^{i,j}, \epsilon_t^{k,l}) - \kappa^{i,j} \kappa^{k,l} \sigma_{i,k}^2 \end{aligned} \quad (19)$$

And again, we can replace $Cov(\epsilon_t^{i,j}, \epsilon_t^{k,l})$ with its sample equivalent to form the relevant moment condition.

4.4 Implementation

Together, equations (15), (16), (17), (18), and (19) form sufficient moment conditions to identify all the parameters of the Kalman filter, F_t, H_t, R_t and Q_t . Of course, one could estimate these parameters by brute force, piling up all the equations into one big GMM estimation. However, the structure of the moment conditions means

that there is an easier way to do this. Equations (15) and (16) are self-contained in that they depend only on $\sigma_i^2, \rho_1^i, \dots, \rho_L^i$. And equations (17) through (19) are just-identified conditional on estimates for $\sigma_i^2, \rho_1^i, \dots, \rho_L^i$. This means that the model can be estimated in two stages. First, use regular GMM applied to the sample versions of equations (15) and (16) to estimate $\sigma_i^2, \rho_1^i, \dots, \rho_L^i$ separately for each i . Then, solve directly for all the other parameters simply by inverting the sample analogues of equations (17) through (19). Since 1) the errors on the moment conditions for equations (15) and (16) are minimized independent of the other parameters, and 2) equations (17) through (19) are solved with zero error, minimizing the moment errors over all parameter simultaneously cannot improve the fit of the moment restrictions. This will yield the exactly the same estimates as one big GMM using all the moment conditions.¹²

This two-stage approach is also much quicker and easier than the obvious alternative: maximum likelihood. Estimating Kalman parameters this way requires computing the likelihood of the model for each candidate parameter value. Moreover, because the Kalman smoother is a sequential estimator, this calculation cannot be parallelized for speed. This is combined with a curse of dimensionality. Because estimated covariances are so important, the number of parameters in the model grows with the square of the number of components and data series, producing a large number of parameters quickly. In the baseline model $L = 1$, $N = 9$, $K = 2$, and $P = 9$, and so there are 60 parameters to optimize over. Estimating even a very simple likelihood function over such a massive parameter space quickly becomes prohibitive.

4.5 Measuring uncertainty via the delta method

A key contribution of this paper is in trying to accurately capture the uncertainty around the estimated monthly series. This requires taking account of two sources of uncertainty: the Kalman prediction error, which is conditional on parameters, and the parameter estimation uncertainty.

More formally, let Ψ be a given parameterization of the Kalman filter, containing

¹²Strictly, this is only true for *unweighted* GMM. If the moment conditions are weighted then they are no longer separable. This risks producing monthly national accounts estimates which are less precise than they could be. However, as I show later in Section D, the gains from efficient GMM are very small, and typically not worth the extra time or complexity. In the baseline, I employ an intermediate approach, using efficient GMM to weight the estimation of the dynamic time series parameters and then unweighted GMM for the remaining conditions.

all the parameters of the model. Then for each i , the output of the Kalman filter is two time series:

The Kalman mean: $\hat{\mathbb{E}}(y_t^i | I_T, \Psi)$

The Kalman prediction error variance: $\hat{\mathbb{V}}(y_t^i | I_T, \Psi) = \mathbb{E} \left(\left(y_t^i - \hat{\mathbb{E}}(y_t^i | I_T, \Psi) \right)^2 | I_T, \Psi \right)$

where the hats on $\hat{\mathbb{E}}$ and $\hat{\mathbb{V}}$ remind us that the Kalman filter only provides *estimates* of the true conditional mean of y_t^i .¹³

Assuming that the model is correctly specified, and denoting by Ψ^* and Ψ_0 the GMM point estimates and the true value of Ψ respectively, then the consistency of Ψ^* implies that the Kalman mean is a consistent estimate of the conditional mean. That is:

$$\Psi^* \rightarrow \Psi_0 \quad \Rightarrow \quad \hat{\mathbb{E}}(y_t^i | I_T, \Psi^*) \rightarrow \mathbb{E}(y_t^i | I_T) \quad (20)$$

However, since the Kalman filter is conditional on the parameters, the variance does not take into account uncertainty over them. As a result, confidence intervals based on $\mathbb{V}(\hat{y}_t^i | I_T, \Psi)$ will be too narrow.

To see this, let $F_n(\Psi)$ be the sampling distribution for Ψ for a sample of size n . Then:

$$\begin{aligned} \text{Var}(y_t | I_T) &= \mathbb{E} \left[(y_t - \mathbb{E}(y_t | I_T))^2 | I_T \right] \\ &= \int \mathbb{E} \left[(y_t - \mathbb{E}(y_t | I_T, \Psi))^2 | I_T, \Psi \right] dF_n(\Psi) \\ &= \int \mathbb{E} \left[\left(y_t - \hat{\mathbb{E}}(y_t | I_T, \Psi^*) \right)^2 | I_T, \Psi \right] dF_n(\Psi) \\ &\quad + \int \mathbb{E} \left[\left(\hat{\mathbb{E}}(y_t | I_T, \Psi^*) - \mathbb{E}(y_t | I_T, \Psi) \right)^2 | I_T, \Psi \right] dF_n(\Psi) \\ &= \underbrace{\int \hat{\mathbb{V}}(y_t | I_T, \Psi) dF_n(\Psi)}_{\text{Avg. Kalman pred. error variance}} + \underbrace{\int \left(\hat{\mathbb{E}}(y_t | I_T, \Psi^*) - \mathbb{E}(y_t | I_T, \Psi) \right)^2 dF_n(\Psi)}_{\text{Variance of the Kalman mean over the parameters}} \end{aligned}$$

where the first line follows from the law of iterated expectations. The final line says that the error variance including parameter uncertainty is the sum of average the

¹³Recall that I_T is the information set containing the data.

Kalman prediction error variance plus the variance of the Kalman mean over the estimated parameters. Given Ψ^* is consistent, the point estimate for the Kalman error variance is a consistent estimate of the first term. To compute the second term, I use the delta method. That is, I compute the sequence of derivative of Kalman means in each period with respect to the parameter vector, Ψ . This is relatively straightforward to calculate numerically by perturbing the parameter vector Ψ in each dimension and computing the marginal change in the resulting Kalman mean.

I thus calculate the covariance for the estimated national accounts components by:

$$\Sigma_t = \underbrace{\hat{V}(y_t|I_T, \Psi^*)}_{\text{Filter uncertainty}} + \underbrace{\frac{d\hat{\mathbb{E}}(y_t|I_T, \Psi^*)'}{d\Psi} \text{Var}\Psi^* \frac{d\hat{\mathbb{E}}(y_t|I_T, \Psi^*)}{d\Psi}}_{\text{Estimation uncertainty}} \quad (21)$$

where Ψ is the just the asymptotic GMM estimator covariance for Ψ . Consistency of Ψ^* means that Σ_t is a consistent estimator of $\text{Var}(y_t|I_T)$

Appendix B summarizes the steps required to produce the final monthly national accounts estimates.

5 Results

5.1 Estimated equations

The first step is to de-trend the data. The results are not essential so details on the estimated polynomial trends are deferred to Table A1 in Appendix A.

Table 3 reports the coefficients of the dynamic state equations. To determine the lag length L_i of each dynamic state equation, I apply the BIC-based consistent model and moment test of Andrews and Lu (2001). This selects a one-lag model for all variables. I also use this test to assess the number of over-identifying restrictions to include, i.e. the number of autocorrelations of the quarterly data to use in equation 16. With $L = 1$, the test fails to reject arbitrarily many over-identifying moments, suggesting that there is little extra information gained by adding more data to the model.¹⁴ Accordingly, I somewhat arbitrarily use eight quarterly autocovariance mo-

¹⁴Intuitively, since the long-lagged auto-correlation of any time series processes tends to the largest characteristic root, the ratio of autocovariances at lag j and $j + 1$ becomes constant as j gets very large. That is, ever more autocovariance moments contain no extra information about the underlying parameters.

ments, but in Section 6.4 I check that using a different number of over-identifying restrictions changes the results little.

The GMM estimates in Table 3 generally produce highly persistent dynamic processes, with monthly persistences for all components except inventories over 0.95. One concern about these estimated processes is that, since they estimate monthly processes using quarterly data, they may be failing to pick up important higher-frequency variation. As a test of this, I also estimate the dynamic state equations for goods and services consumption directly from the monthly data by ordinary least squares (columns headed “OLS”) in Table 3. Although the coefficients are not exactly the same (since the use different information sets – GMM matches nine moments and OLS implicitly only two), they produce similar estimates, with a difference in the estimates for ρ_1 of less than one standard error. The fit of the equations is also good, with the average moment error less than 2 percent for all NIPA components except the volatile investment series.

	OLS		GMM							
	Services	Goods	Services	Goods	NRFI	RFI	Gov.	Exports	Imports	Inv.
ρ_1	0.973*** (0.0082)	0.969*** (0.00878)	0.979*** (0.00421)	0.978*** (0.00361)	0.965*** (0.00527)	0.971*** (0.00584)	0.96*** (0.00954)	0.976*** (0.00399)	0.985*** (0.003)	0.808*** (0.0334)
$100 \times \sigma_{i,i}^2$	0.00348*** (0.000285)	0.0163*** (0.00133)	0.00267*** (0.000716)	0.0104*** (0.00192)	0.0419*** (0.00691)	0.161*** (0.0355)	0.0282*** (0.0091)	0.0608*** (0.0126)	0.0479*** (0.0108)	66.8*** (17.3)
Observations	795	795	301	301	301	301	301	301	301	301
Moment error			0.001	0.015	0.055	0.049	0.007	0.005	0.004	0.016

Table 3: Estimated coefficients for the state equations

Table reports the estimated coefficients for the state equations. Columns labeled “GMM” use quarterly data to estimate coefficients by GMM as described in Section 4.2, with asymptotic standard errors in parentheses. Those labeled “OLS” are for comparison, and use monthly data for the two exact series. All GMM estimates use eight lags of the quarterly data, i.e. $k = 1, \dots, 8$ in equation (16). The row labeled “Moment error” is average weighted square error on the moment conditions, in percent deviation from trend. Sample: 1950Q1-2025Q1 (GMM), Jan 1959-Mar 2025 (OLS).

As mentioned before, the cross-correlation of the states is an important input into the Kalman filter. The joint movement of the components is defines the extent to which information about one component affects the estimates of the others. Figure 3 reports the estimated innovation correlation matrix for the dynamic states. Some patterns stand out. The four components of domestic private demand, the investment and consumption series, are highly correlated, consistent with the notion that demand shocks drive much of the variation in these series. Likewise, the importance

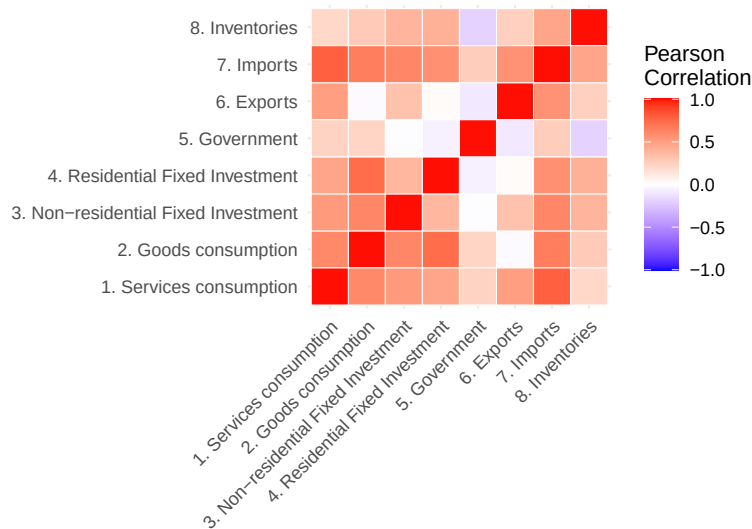


Figure 3: Estimated correlation of innovations to hidden states.

Figure shows correlations of the innovations to the monthly GDP components, estimated by GMM. Colors in each cell correspond to the value for $\sigma_{i,j}^2 / (\sigma_{i,i} \sigma_{j,j})$.

of imported goods in consumption and investment explains the high correlation of these series (especially goods consumption) with imports. In contrast, exports and government spending are much less correlated with domestic demand components.

Table 4 reports the estimated coefficients for the $\kappa^{i,j}$ coefficients for the non-observed NIPA components. In general, coefficients are highly significant and the statistical fit of the estimated relationships is generally strong.¹⁵ In the case inventories, there is a trade-off between the informativeness of the monthly indicator, and the sample for which the indicator is available – the manufacturing inventories series is a worse proxy for actual inventories than monthly real inventories, but is available for a shorter period.

5.2 Estimated Monthly GDP

Figure 6 shows the estimated monthly series for GDP and its components for the full sample. Since the details of the series are hard to see over long periods of time, Figures 4 and 5 present more detailed views of the estimated series near the start and end of the sample, including confidence intervals computed using the variance calculation in

¹⁵I report a pseudo- R^2 because, unlike OLS, GMM estimates are not bound to lie inside $[0, 1]$.

y_t^i	$q_t^{i,j}$	$\kappa^{i,j}$	$Cor(E_t^{i,j}, \hat{Y}_t^i)$	N_{obs}	Pseudo- R^2
Exports	Census Exports Volumes	0.595*** (0.1)	0.21	133	0.62
Imports	Census Imports Volumes	-0.0835 (0.213)	-0.02	133	0.01
Inventories	Change in Real Inventories	1.26*** (0.134)	0.75	113	0.90
	Change in manufacturing inventories	0.628*** (0.198)	0.23	268	0.12
Non-residential Fixed Investment	Non-residential construction spending	0.313*** (0.0893)	0.22	93	0.26
Residential Fixed Investment	Housing starts	0.416 (0.332)	0.07	261	0.08
	Housing starts, total multifamily	0.511 (0.765)	0.03	265	0.02

Table 4: Estimated coefficients for the observation equations

Table reports the estimated coefficients for the observations equations. Estimates are by GMM, using method outlined in Section 4.2. Standard errors computed by asymptotic GMM. Data sample is quarterly and varies with $q_t^{i,j}$, starting in the first complete quarter given monthly start dates in Table 2 and ends in 2025Q1. Column “ N_{obs} ” gives the number of quarterly observations for each estimated relationship. The Pseudo- R^2 is a measure of the implied fit of the monthly observation equation, computed as the ratio of explained to total variance: $(\kappa^{i,j})^2 \text{var } w_t^i / \text{var } \epsilon_t^{i,j}$. Finally, the column labeled “ $Cor(E_t^{i,j}, \hat{Y}_t^i)$ ” is the sample correlation of the quarterly data and is included only for reference. Sample: various-2025Q1

Section 4.5.¹⁶ In Figure 4, the confidence intervals are relatively wide since no relevant monthly indicators are available prior to 1959. Consistent with this, the estimated monthly time series is much smoother. However, after 1959 things change with the introduction of the monthly consumption data. Thereafter, the estimates are more informed and so become more volatile. Note in particular how much more sharply-identified are the declines in the investment series in the 1960 recession versus the one in 1957. The correlation of state residuals also means that confidence intervals on all variables shrink, not just on consumption (more on this below). By the onset of COVID (Figure 5) much more monthly indicators are available, resulting in much more precise series and tighter confidence intervals. The abrupt halt in activity due to COVID-19 is precisely identified as staring in March 2020, as one would hope.

¹⁶Charts for four other episodes are presented in Appendix C.

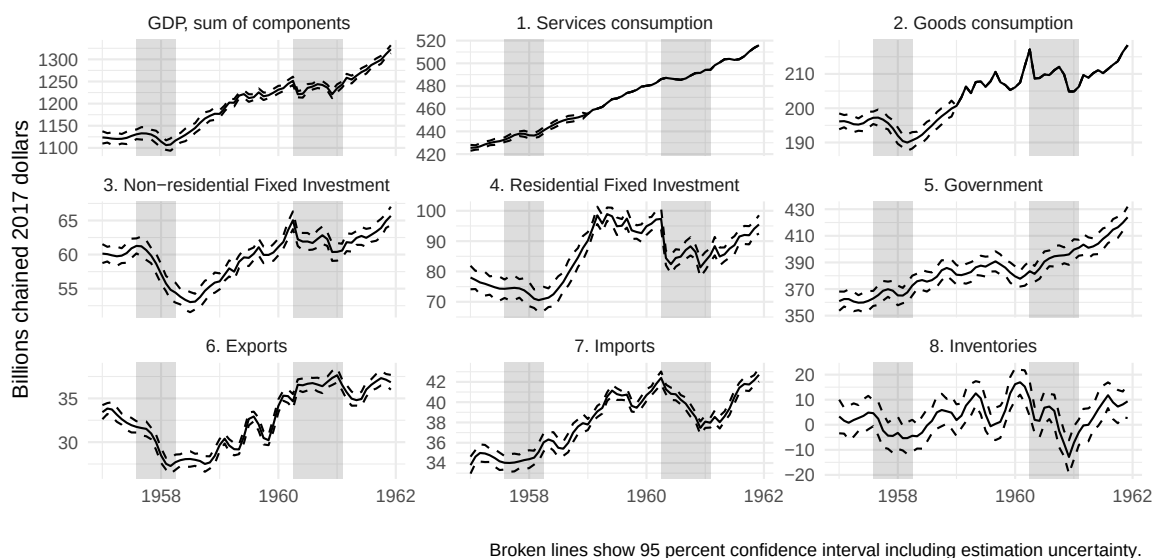


Figure 4: Estimated Monthly National Accounts Series: Late 1950s

Figure shows estimated monthly national accounts series plus 95 percent confidence interval. Shaded areas are NBER recessions. Sample: Jan 1957-Dec 1961.

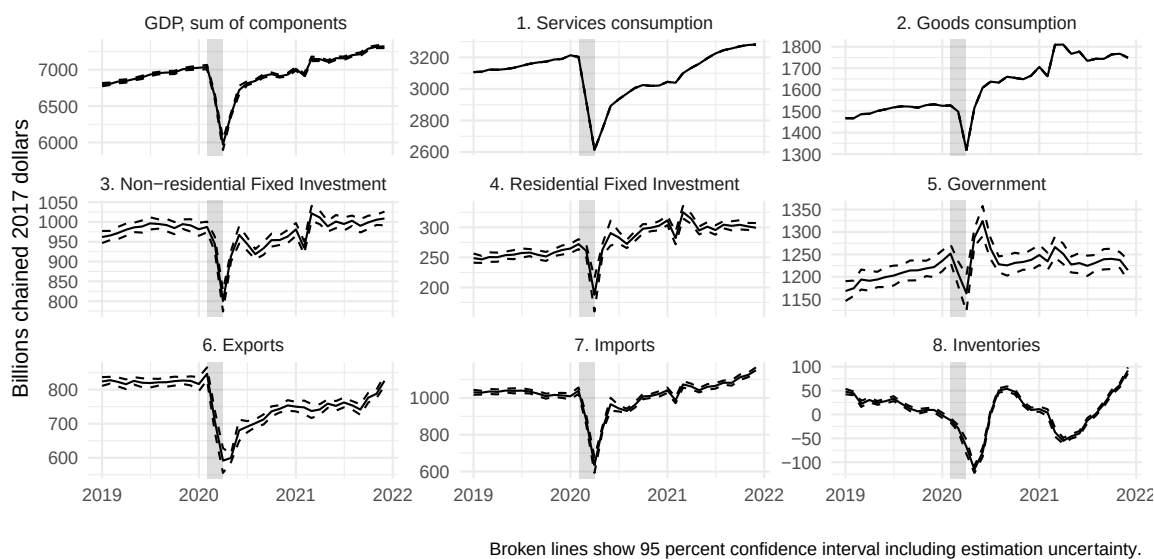


Figure 5: Estimated Monthly National Accounts Series: COVID-19

Figure shows estimated monthly national accounts series plus 95 percent confidence interval. Shaded areas are NBER recessions. Sample: Jan 2019-Dec 2022.

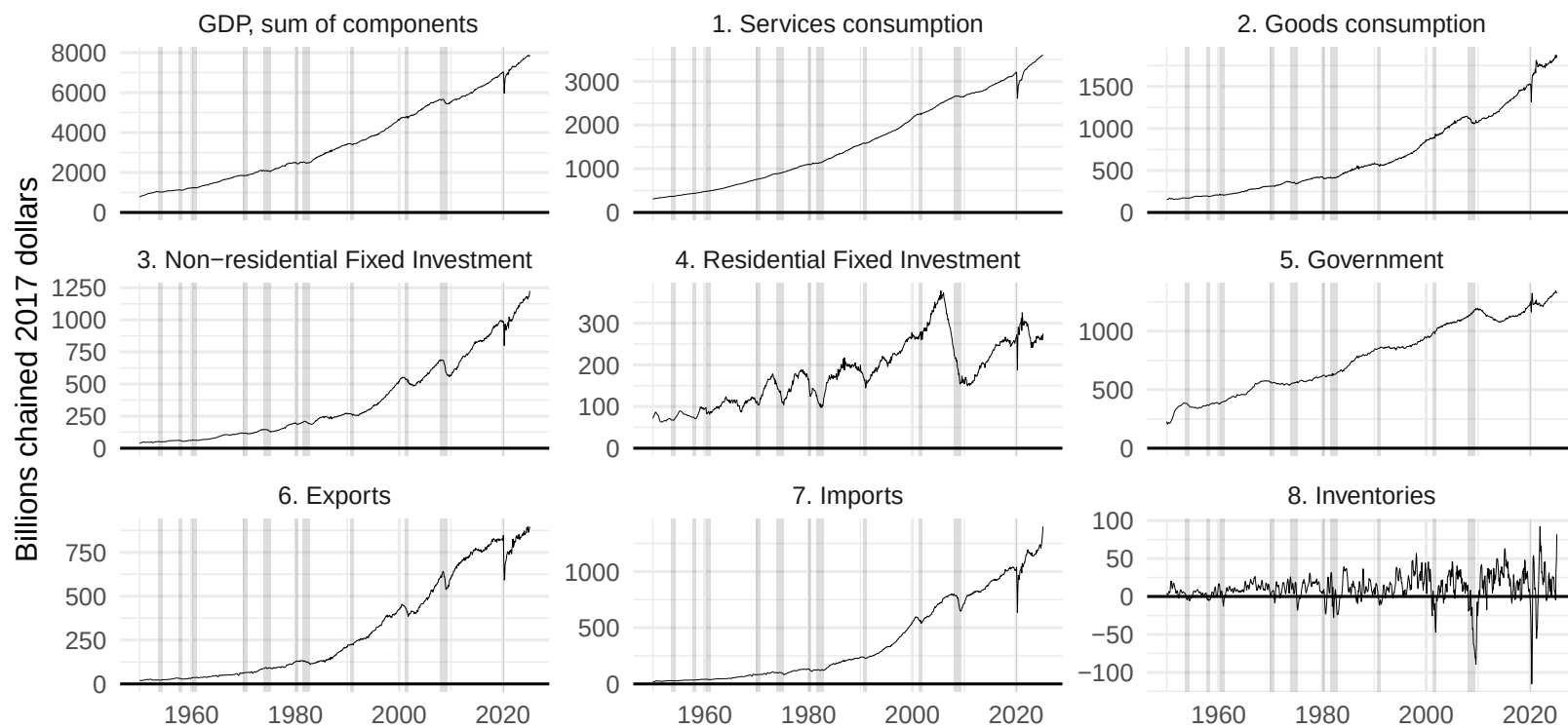


Figure 6: Estimated Monthly National Accounts Series: Point Estimates

Figure shows estimated monthly national accounts series plus 95 percent confidence interval. Shaded areas are NBER recessions. Sample: Jan 1950-Mar 2025.

5.3 Measuring uncertainty

The confidence intervals for the estimated monthly series are narrow. As shown in Figure 7, the standard errors for GDP are 0.3 percent of trend or less since the 1960s, and less than 0.2 percent since the late 1990s (COVID is a notable exception). The gradual decrease over time is principally due to the increasing share of consumption (which is known precisely) in GDP. However, the introduction of monthly data on exports and inventories in the 1990s dramatically reduces the uncertainty stemming from these components.

Figure 7 also shows the information spillovers between components. Most notably in 1959, monthly consumption data become available. This has a direct and obvious impact on the standard errors for the two consumption series, reducing them to zero. But the uncertainty for other components – notably the two investment series and imports – also declines, by almost one half in the case of imports. Note that there is no additional information directly relevant to these series added to the model at this point, the only extra information is the two consumption series. The reduction in uncertainty for the investment and imports series occurs as a result of the strong estimated correlation residual correlation between these series and consumption shown in Figure 3. Intuitively, because investment and imports are positively correlated with consumption, more precise signals about consumption are informative for investment and imports.

Finally, Figure 7 also presents the decomposition of uncertainty over monthly series into the two components in equation (21). Most of the time, the main source of uncertainty over the monthly estimates comes from the filter. This is constant over intervals where there is no change in the information set available to the model (except for GDP, which is calculated a sum of components, and so the filter uncertainty changes as the shares of those components change). In contrast, estimation uncertainty is not constant, and varies over time. Mechanically, this arises because estimation uncertainty enters multiplicatively with the data. The estimates are linear combination of the data, where the coefficients are (combinations of) the estimated parameters. And because these coefficients are multipliers of the data, the data are multipliers of the uncertainty on the coefficients. This is why uncertainty over the true monthly series increases in times of volatility. This figure also makes clear the tradeoff between filter and estimation uncertainty in model selection. For example, in January 1992 the inclusion of monthly census data on exports adds an extra piece of

information to the model. Conditional on knowing the correct relationship between this data and true monthly NIPA-equivalent exports, this mechanically reduces the uncertainty; more data means less uncertainty. This is seen in the drop in the filter uncertainty at this time. But this relationship is not known for certain. Thus adding extra data comes at the price of including an additional uncertain parameter. This is why higher estimation uncertainty partially offsets the reduction in filter uncertainty when monthly export data become available. This highlights the model selection problem: adding more data reduces filter uncertainty but by adding parameters it increases estimation uncertainty.

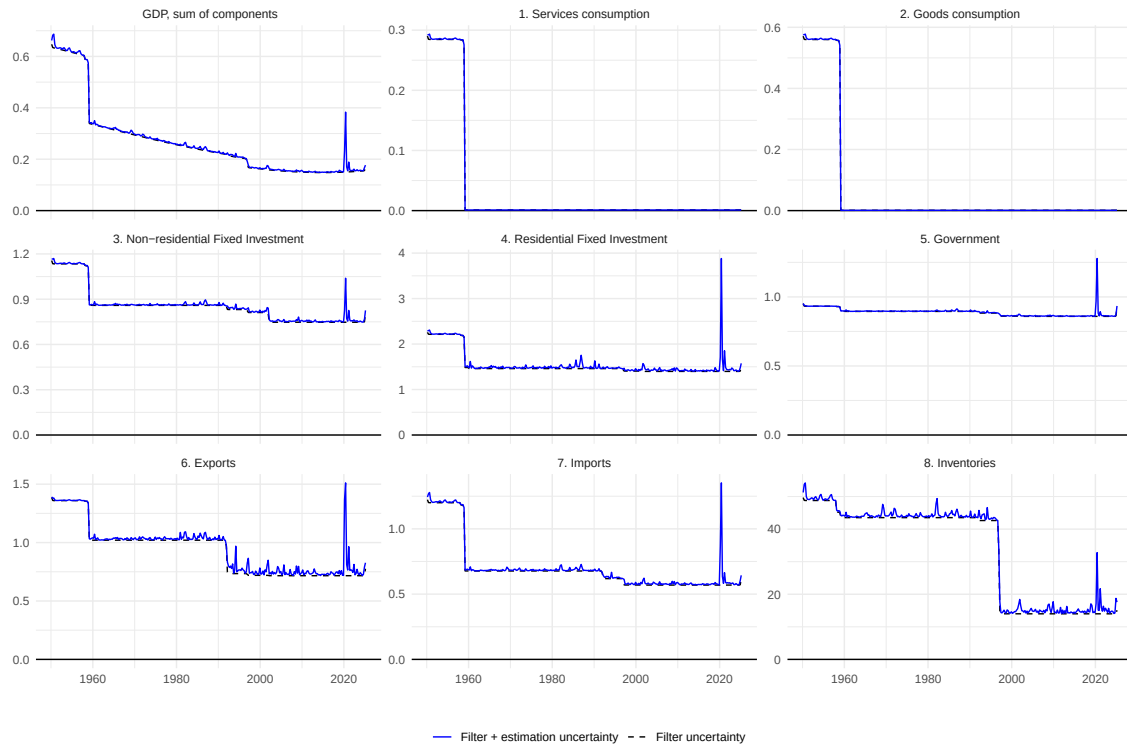


Figure 7: Times series standard deviation by component

Figure shows quarterly average standard error of the estimated monthly national accounts series, as a percentage of trend. Quarterly averaging suppresses systematic within-quarter volatility due to periodic NIPA data releases. Sample: 1950Q1-2025Q1.

A summary measure of the utility of the estimates and their uncertainty is given by the signal-to-noise ratio. This is typically defined as the ratio of the variance of a noisy estimate (the signal) to the variance of noise (the difference between the the signal and the true realization). The idea is that variation in the signal is good but

	$\sqrt{Var\mu_t^{gdp}}$	$\sqrt{Var_t\mu_{t+1}^{gdp}}$	$\sqrt{\Sigma_t^{gdp,gdp}}$	Ratios	
Decade	(1)	(2)	(3)	(1)/(3)	(2)/(3)
1950	3.87	1.32	0.56	6.9	2.3
1960	3.23	1.20	0.30	10.7	3.9
1970	2.17	1.16	0.27	8.1	4.4
1980	2.81	1.16	0.24	11.5	4.7
1990	1.86	1.17	0.20	9.4	5.9
2000	2.43	1.22	0.15	16.0	8.0
2010	0.72	1.30	0.15	4.7	8.5
2020	2.79	1.40	0.18	15.6	7.8

Table 5: Signal-to-noise ratios: GDP by decade

Table shows averages by decade of the standard deviation of the (1) unconditional and (2) one-step-ahead forecasts of monthly GDP, and well as the model-implied standard error (3). Units are in percentage of trend. Sample: Jan 1950-Mar 2025.

only to the extent that it is informative. One difficulty in implementing this in the current context is that the the NIPA series have trends, and so we cannot just use the variance of the point estimate for GDP as the signal variance, since it will be a function of the trends as well. One way to address this is to use the variance of the detrended GDP estimate as the signal variance, $Var\mu_t^{gdp}$. However, this is not a perfect solution since it will still include some predictable business cycle fluctuations in the signal variance, which are not really a measure of the informativeness of the monthly estimate. So my preferred measure of the signal variance is the one-step-ahead prediction variance of detrended GDP, $Var_{t-1}\mu_t^{gdp}$. This is the variance of the detrended GDP series in period t conditional on information in period $t - 1$, and is approximately the variation in one-period-ahead growth rates.

Table 5 reports these two measures of the signal-to-noise ratio for GDP by decade. The preferred signal-to-noise ratio is in the rightmost column. This says that the one-month-ahead variation in the estimated monthly GDP series is about 2 times the standard error of the GDP in 1950s and 7 times in the 2010s (and about 4.5 times on average over the whole sample). Put differently, the monthly GDP estimate in the 2010s is precise enough that its standard error is only one seventh of the volatility of the one-month growth rate of GDP. Table A2 in Appendix A presents the results by component. It finds that, except for consumption where we have exact data for

most of the sample, the signal-to-noise ratio is highest for imports and lowest for inventories.

5.4 Recession Dating

Via its business cycle dating committee (BCDC), the NBER produces monthly start and end dates of US recessions. The BCDC defines a recession as “a significant decline in economic activity that is spread across the economy and that lasts more than a few months”, and identifies them using a broad set of indicators with considerable latitude on the importance placed on each. Given its importance as an aggregate summary measure of economic activity, a monthly GDP data would likely be a key input into this process.¹⁷

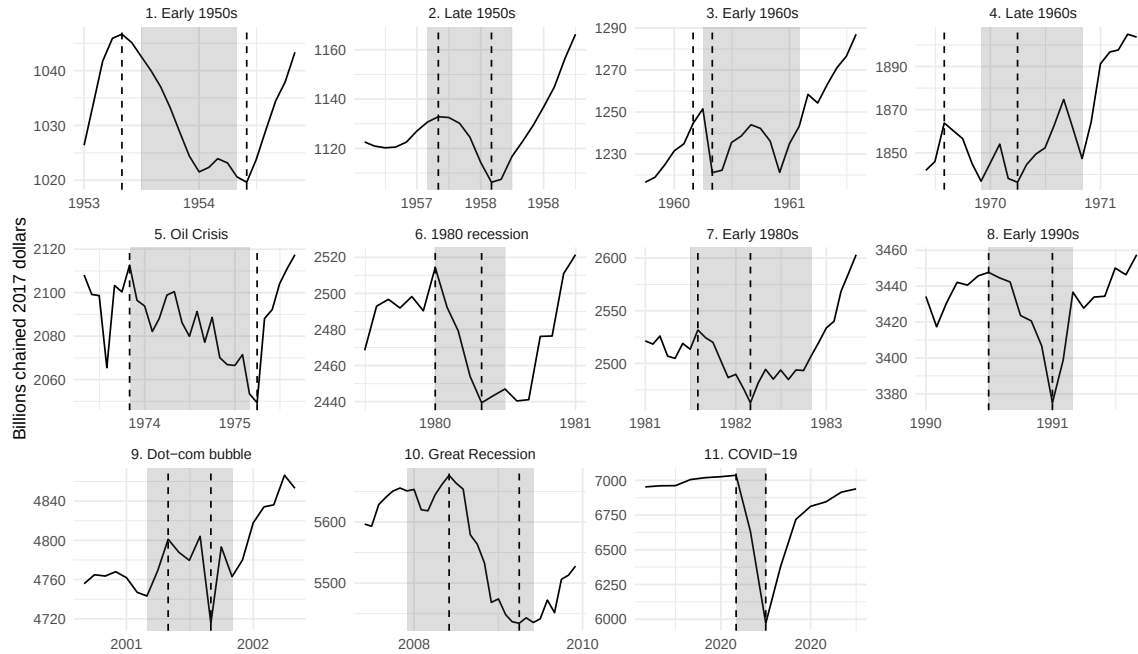


Figure 8: Recession Dates: NBER versus monthly GDP

Figure shows estimated monthly GDP around 11 recessions identified by the NBER. Shaded areas show NBER recessions and vertical lines local peaks and troughs in estimated monthly GDP. Sample: Jan 1960 - Dec 2019.

¹⁷A point made implicitly by the NBER itself when discussing quarterly and monthly dating, saying “Two measures that are very important in the determination of quarterly peaks and troughs, but that are not available monthly, are [...] GDP and GDI.” See: <https://www.nber.org/research/business-cycle-dating/business-cycle-dating-procedure-frequently-asked-questions>.

Recession	NBER		Monthly GDP	
	Peak	Trough	Peak	Trough
1. Early 1950s	Jul 1953	May 1954	May 1953 (Mar 1953,Aug 1953)	Jun 1954 (Nov 1953,Aug 1954)
2. Late 1950s	Aug 1957	Apr 1958	Sep 1957 (Dec 1956,Dec 1957)	Feb 1958 (Jan 1958,Apr 1958)
3. Early 1960s	Apr 1960	Feb 1961	Mar 1960 (Mar 1960,Apr 1960)	May 1960 (May 1960,Dec 1960)
4. Late 1960s	Dec 1969	Nov 1970	Aug 1969 (Aug 1969,Oct 1969)	Apr 1970 (Nov 1969,May 1970)
5. Oil Crisis	Nov 1973	Mar 1975	Nov 1973 (May 1973,Nov 1973)	Apr 1975 (Mar 1975,Apr 1975)
6. 1980 recession	Jan 1980	Jul 1980	Jan 1980 (Jan 1980,Jan 1980)	May 1980 (May 1980,Sep 1980)
7. Early 1980s	Jul 1981	Nov 1982	Aug 1981 (Jan 1981,Sep 1981)	Mar 1982 (Feb 1982,Mar 1982)
8. Early 1990s	Jul 1990	Mar 1991	Jul 1990 (Jan 1990,Sep 1990)	Jan 1991 (Jan 1991,Jan 1991)
9. Dot-com bubble	Mar 2001	Nov 2001	May 2001 (May 2001,Aug 2001)	Sep 2001 (Sep 2001,Sep 2001)
10. Great Recession	Dec 2007	Jun 2009	Jun 2008 (May 2008,Jul 2008)	Apr 2009 (Feb 2009,Sep 2009)
11. COVID-19	Feb 2020	Apr 2020	Feb 2020 (Dec 2019,Feb 2020)	Apr 2020 (Apr 2020,Apr 2020)

Table 6: Recession dates: NBER versus monthly GDP

Columns show dates of local peaks and troughs of two versions of real activity: that used by the NBER recession dating committee and estimated monthly GDP. Monthly GDP series include 95 percent confidence intervals in parentheses, computed as the range of dates where point estimates are within two standard deviations of the estimated peak or trough. NBER dates are in bold when they fall outside the equivalent monthly GDP confidence intervals. Sample: Jan 1950-Mar 2025.

To see how monthly GDP data might paint a different picture of US economic history, I compare the NBER recession dates since 1950 to an alternative monthly recession dating approach using just estimated monthly GDP. To identify recessions, I identify a series of peak-trough pairs, where each peak is the highest level of GDP prior to its paired trough, and the trough is the lowest level of GDP after its paired peak. Imposing that peaks and troughs cannot be in consecutive months, I define pairs satisfying:

$$peak_j = \arg \max_{t < trough_j - 1} gdp_t \quad \quad \quad trough_j = \arg \min_{t > peak_j + 1} gdp_t \quad (22)$$

where j indexes the pairs.

Figure 8 presents the results of this exercise, compared to NBER recessions in the same period. In general, the start dates of the recessions are fairly closely aligned, but end dates are more often earlier than the NBER dates. Of course, some of the difference in the recession dating arises because the monthly GDP estimates are uncertain. And so Table 6 adds 95 percent confidence intervals to the recessions dates. This computes the range of dates where the point estimate for monthly GDP is within the 95 percent confidence interval at the peak (or trough).¹⁸ Of the eleven recessions, start dates for three – the late 1960s, the dot-com bubble in 2001, and the Great Recession – are statistically different. As can easily be seen in Figure 8, all three episodes feature stuttering starts to the downturn, in contrast to sudden onset recessions, such as the early 1960s or the COVID-19 pandemic. In such cases, there is more than one reasonable choice for the start date. For example, one could equally date the start of the Great Recession to either late 2007, when initial signs of stress began to appear, or to late Summer 2008, just before the financial sector collapsed. The dating of the end of business cycles, however, shows larger and more systematic differences between the two methods. Five of the eleven recessions have statistically significantly different end dates across the two methods, usually with earlier end dates using monthly GDP. In contrast to the dating of peaks, where more than one starting point seems intuitively reasonable, the troughs are almost all cases clear low points in the monthly GDP series. One interpretation is that the BCDC is slightly conservative in calling recessions, waiting for a pick-up in a broad swathe of indicators, including some, such as labor market indicators, which may lag real activity.

5.5 Monthly Business Cycle statistics

To see how the estimated monthly national accounts would give a different picture of business cycle dynamics than the quarterly data, Table 7 presents a variance decomposition for business-cycle fluctuations in GDP by component for each frequency.

The decomposition uses the formula:

$$var \tilde{gdp}_t \simeq \sum_{i=t}^N \omega_t^i cov(\tilde{y}_t^i, \tilde{gdp}_t)$$

¹⁸An example: In the early 1960s, the point estimate for monthly GDP peaks in April 1960, but point estimates for March to October 1960 are all higher than the 95 percent lower bound for April.

where $\tilde{y}_t^i$ and $\tilde{gdp}_t$ are the business-cycle-frequency filtered variables, and $\omega_t^i = y_t^i/gdp_t$ is the level GDP share.¹⁹

The main point that stands out from Table 7 is that the quarterly data understate the importance of government spending in changes in GDP early in the sample. Through until the 1980s, the fraction of variation in monthly-frequency GDP due to government spending is much higher in the monthly estimates, suggesting that quarterly data might underestimate the shift away from fiscal demand management.

Of course, changes in variance contributions could be due to different cyclical correlations, or a combination of weights and relative variances. And so Table A3 in Appendix A repeats the same exercise for the more familiar correlations of components with GDP. Generally, the correlation of all subcomponents with GDP is lower at the monthly frequency, with non-residential fixed investment showing a much weaker correlation. This suggests that in general, there is component-specific information not correlated with overall demand which is driving these series, although that this is largely offset by changes in relative variances and GDP share in the variance decomposition in Table 7. The exception is government spending, where the correlation is much higher in the monthly data, consistent with changes in systematic fiscal policy driving the changing importance of government spending in business cycles.

6 Validation and Robustness

6.1 Comparison to quarterly NIPA data

To check that the accounting constraints hold as they should, Figure 9 compares the summed monthly data to the quarterly NIPA statistics. There is essentially no discrepancy for any of the components. The one exception is for GDP, where the aggregation error in the NIPA data causes a small discrepancy early in the sample. Figures A2 and A3 in Appendix A plot the error relative to the NIPA data for the model with and without the aggregation error and shows that this mitigates the discrepancy.

¹⁹The relationship is approximate because the business cycle filtering process does not respect additivity. Table 7 thus includes an error component.

Decade	Services	Goods	NRFI	RFI	Gov.	Exports	Imports	Inv.	Err.
<i>Quarterly</i>									
1960s	14.8	20.6	13.2	3.6	34.3	-1.4	-5.8	9.6	11.0
1970s	12.5	22.2	15.2	31.4	-1.5	3.6	-12.1	13.9	14.9
1980s	20.2	15.5	7.6	22.8	20.4	7.8	-15.1	17.4	3.6
1990s	20.2	26.2	34.8	21.4	-8.8	13.4	-34.1	11.5	15.4
2000s	18.3	28.5	33.5	29.0	-11.1	21.7	-45.2	23.7	1.6
2010s	36.9	17.0	24.2	10.8	9.0	-10.1	0.9	13.8	-2.5
2020s	67.8	16.7	13.9	3.2	-2.9	19.5	-32.4	13.4	0.8
All	25.7	19.7	18.9	16.3	8.2	7.5	-18.5	14.6	7.5
<i>Monthly</i>									
1960s	16.0	24.3	14.2	8.0	43.8	-2.7	-6.6	11.5	-8.5
1970s	12.9	25.0	12.6	36.2	10.9	1.0	-14.9	12.8	3.4
1980s	22.6	18.5	4.5	28.5	25.7	4.5	-16.3	15.9	-3.8
1990s	25.2	32.8	39.6	25.4	-5.0	13.3	-38.4	15.1	-8.0
2000s	17.6	28.4	31.8	30.1	-10.3	19.4	-44.0	23.6	3.4
2010s	38.7	18.1	21.8	8.4	13.6	-6.4	-4.8	14.3	-3.5
2020s	62.4	20.8	15.6	4.9	1.8	18.8	-34.9	10.0	0.4
All	25.6	21.4	18.0	17.3	19.3	6.4	-20.2	13.7	-1.5

Table 7: Business-cycle statistics: GDP variance shares

Table shows percent of variance in GDP attributable to each component at business cycle frequencies. See text for details of decomposition. All series are first filtered using a Christiano-Fitzgerald filter to eliminate fluctuations at horizons longer than 5 years. Monthly data use point estimates. “Services” and “Goods” columns are private consumption of services and goods respectively. “Err.” is the approximation error in the variance decomposition. Sample: Jan 1960-Mar 2025.

6.2 Testing against known data

In Section 5.3 I argued that the confidence intervals for the estimated series are tight. But are they correct? That is, would the true value of GDP, if it were known, fall inside a x percent confidence interval x percent of the time? Given that monthly GDP data do not exist, it is impossible to provide a direct answer to this question. Instead, I provide indirect evidence that the confidence intervals are valid.

To do this, I re-estimate the model without the monthly goods consumption data. That is, I pretend that the x_t^i “hard” data series for goods consumption does not exist and treat this series as an unobserved state to be inferred from noisy indicators, just like any of the other GDP components, re-estimating the model accordingly. This produces a monthly sequence of confidence intervals for goods consumption. However,

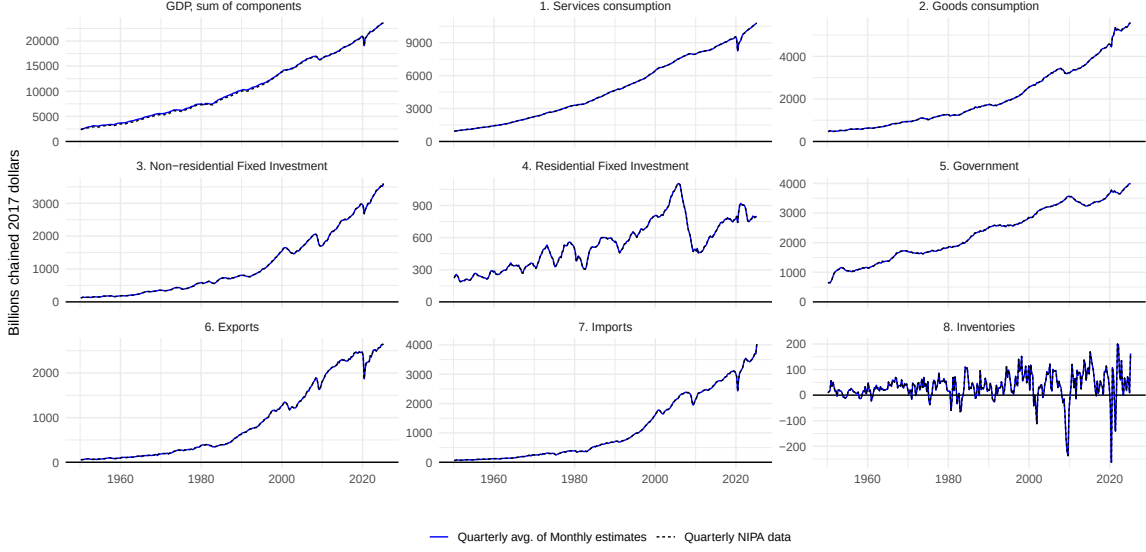


Figure 9: Comparison with quarterly NIPA data

Figure shows quarterly sum of monthly NIPA estimates with actual quarterly NIPA data. Sample: 1950Q1-2025Q1.

since we do have data for goods consumption, we can evaluate the coverage ratios of the confidence intervals for goods consumption. I do this by computing the p-value of the outturn under the distribution implied by the baseline model. That is, I compute:

$$p_t^{goods} = \Phi \left(\frac{x_t^{goods} - \mu_t^{goods}}{\sqrt{\Sigma_t^{goods,goods}}} \right) \quad (23)$$

Where μ_t^{goods} and $\Sigma_t^{goods,goods}$ are the mean and variance of the monthly goods consumption estimates. This calculation is useful because if model uncertainty correctly reflects the true distribution of outcomes, then the distribution of p-values should be uniform.²⁰

Figure 10 plots the cumulative distribution of p-values for two versions of the model with different noisy monthly indicators for goods consumption. One uses only durable goods consumption data (a subset of total goods consumption) and the other real retail sales.²¹ A perfect model would produce plots where the empirical p-values

²⁰Computing p-values is also more general than computing coverage ratios of given confidence intervals, since it characterizes the accuracy of the entire distribution, not just in certain ranges.

²¹I omit the 1950s and the 2020s from this plot. The former is excluded because it is identical to

from the model line up identically along the 45 degree line. Both models are remarkably close to this, suggesting that the confidence intervals at least for this version of the model are remarkably good. This agreement is not pre-ordained. Throughout the sample, different data series become available at different times (not only real retail sales, which starts in 1967, but also spillovers from the other series) causing the appropriate degree of uncertainty to change in response. Table 8 formalizes this comparison, reporting the p-values for the Kolmogorov-Smirnov test that the distributions in Figure 10 are uniform, with the durable goods model doing well throughout.

This exercise also highlights the importance of capturing both filter and estimation uncertainty. The confidence intervals used in Figure 10 use both, but Figure 11 repeats the exercise using only the filter uncertainty. In this case, the resulting p-values are far from correct.

The foregoing exercise is convincing evidence that a version of the model without hard data on goods consumption can correctly characterize uncertainty for that individual series. But what does this imply for the performance of the baseline model? This is a different model, with a different data set and results do not map over identically. This test does at least mimic how the model works in the baseline; all the non-consumption data remain exactly the same and the estimation process is identical to the main exercise in this paper. So even if the results for this special case do not necessarily carry over to the baseline model, this validation exercise is at least an opportunity for the model to fail. And if it had, one would almost certainly conclude that the rest of the model performs poorly.

Consumption data	1960s	1970s	1980s	1990s	2000s	2010s	All
Using durable goods data	0.078	0.334	0.444	0.971	0.215	0.857	0.145
Using real retail sales	0.000	0.062	0.016	0.016	0.004	0.982	0.000
No monthly data	0.036	0.104	0.110	0.095	0.040	0.955	0.000

Table 8: Kolmogorov-Smirnov p-values for validation using goods consumption

Table reports the p-values for the Kolmogorov-Smirnov test that the distribution of forecasts for goods consumption equals that in the data. All models are re-estimated without the monthly goods consumption data. Lower p-values are stronger evidence that the forecast and outturn distributions differ. Column titled “Consumption data” describes what goods-consumption-relevant monthly data the model does see. Sample: Jan 1950-Dec 2019.

the baseline model before 1959, since there is no relevant monthly data. And the latter is dropped because there is only half a decade to test against, and so the results are not comparable.

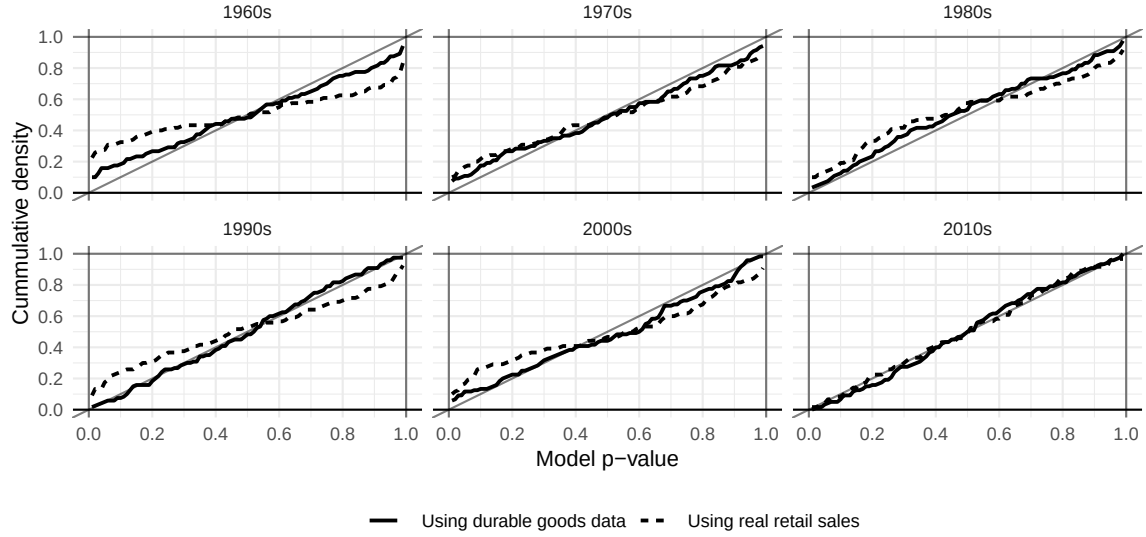


Figure 10: Model validation: Goods consumption

Figure shows the cumulative distribution of p-values goods consumption data according to the monthly estimated distribution including filter and estimation uncertainty. All models are re-estimated without the monthly goods consumption data. Perfectly accurate confidence intervals would have p-values exactly on the 45 degree line. Sample: Jan 1960 - Dec 2019.

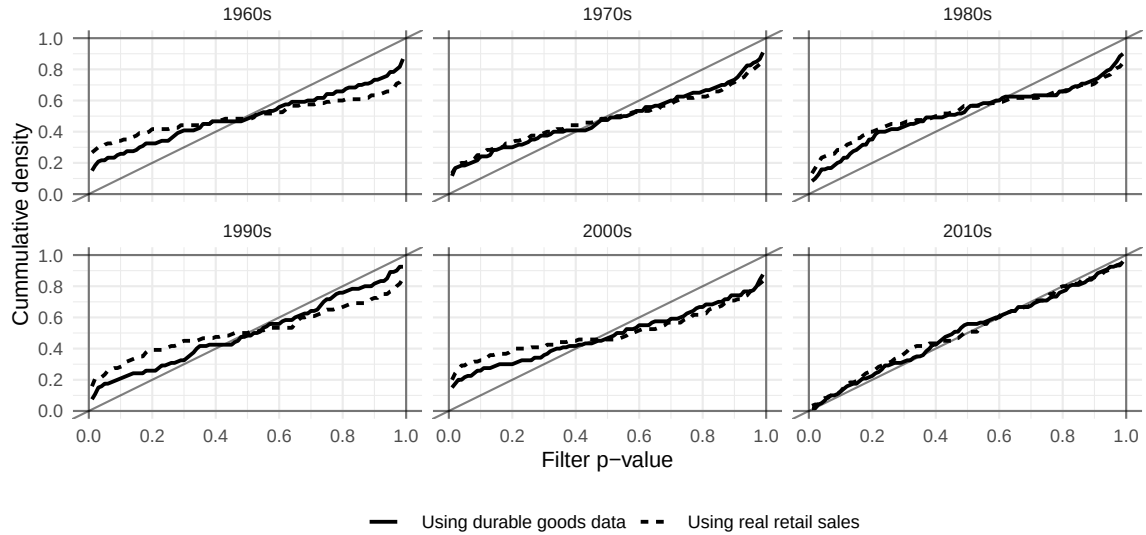


Figure 11: Model validation: Goods consumption, filter uncertainty only

Figure shows the cumulative distribution of p-values goods consumption data according to the monthly estimated distribution using filter uncertainty only. All models are re-estimated without the monthly goods consumption data. Perfectly accurate confidence intervals would have p-values exactly on the 45 degree line. Sample: Jan 1960 - Dec 2019.

6.3 Model Selection

In this section, I address the question of whether there are other data sources which could help improve the accuracy of the estimated monthly national accounts. In doing so, I consider six alternative models, each including extra data on other monthly series, which could be used as measures of activity for particular components. These include various public and sector-specific labor market measures (public wages & salaries, residential construction employees), as well as sector-specific output (nondefense capital shipments), and cost measures (oil prices).

Table 9 reports the results of this exercise, showing in the top half of the table the alternate specifications, and in the bottom half various measures of model performance. In-sample measures of fit, as captured by AIC and BIC suggest that there are gains from adding extra data series to the baseline model. However, the practical improvements of these changes are minuscule. The confidence intervals for GDP in the alternate models are narrower than those in the baseline by at most one one-hundredth of a percentage point across all specifications.

Why does extra data not improve the monthly estimates by more? One reason is simply that the series in the baseline model are already among the best available indicators for monthly activity in the relevant components. Indeed, I include no extra indicators for the trade or inventories series simply because those already included are so close to the data. As such, the data sources remaining for inclusion in alternative models have generally poor correlations with the data (see ‘pseudo- R^2 ’ column in Table 9, which reports the monthly pseudo- R^2 of the bivariate regression, equivalent to that in Table 4). And even when the bivariate correlation is strong, the marginal correlation conditional on other observables can be limited (in particular, the impact of including residential building construction employees). A second reason is that extra variables reduces filter uncertainty but increases estimation uncertainty. This is particularly true when the correlations of observables and NIPA series are weak, and so the estimated relationship is subject to increased uncertainty. The baseline model is large enough, and the extra data poor enough indicators, that, at the margin, the filter and estimation uncertainty effects roughly offset. And yet a third reason is that the implicit criteria for a series to be informative in the Kalman framework is actually quite restrictive. The data have to say something reliable about the underlying series in the current month. If there are long and variable lags between the unobserved NIPA equivalent series and the observable data then this relationship will be very

weak.

Nevertheless, Table 9 shows that there at least some gains from including extra data relative to the baseline. So why not use one of these as the baseline? The answer is that the performance of the models in Table 9 is essentially a tie. In any practical sense, the gains in reducing uncertainty are indistinguishable from zero. In which case, other concerns become important, such as numerical stability. And since the aim of this project is to produce not just a one-time series for monthly GDP, but something which can be routinely updated by readers with the associated code, this matters. Not reported are the (many!) failed runs where including too many near-collinear observable data series causes a matrix inversion to fail.

6.4 Robustness to model specification

To check that the baseline results are not unduly sensitive to assumptions about details of the model specification other than the monthly indicators, I compute estimated monthly national accounts series for in four further cases. In one, I use efficient GMM, rather than unweighted. In another, I fit the model to more dynamic moments, matching 12 autoregressive equations. In a third I consider a higher-order polynomial trend. And in the last specification, I model the aggregation error in the national account specifically as a ninth component of GDP. Tables A4 and A4 in Appendix D report the average and average absolute deviations from the baseline. In all cases except the last, the differences are minimal. When including the aggregation error, difference for the GDP series are naturally larger – GDP itself in the quarterly data is now different, and includes an explicit aggregation error. Reassuringly, though, the inclusion of the aggregation error has no noticeable effect on any of the other components of GDP. In other words, including the aggregation error in the monthly series has exactly the same effect as it does on the quarterly data – it drives a wedge between the components and the total. In the data associated with this paper I report two series for GDP: one including the aggregation error and one excluding it. This means that the choice facing prospective users of the monthly data is identical to that they would face with the quarterly data anyway: either include the aggregation error and match headline GDP exactly, or drop it and have the components sum exactly.

		pseudo- R^2	Baseline	Gov. (small)	Gov. (large)	NRFI	RFI	Everything
<i>Extra variables</i>								
Government	Federal outlays	0.053	✗	✗	✓	✗	✗	✓
	Value of Public Construction	0.001	✗	✗	✓	✗	✗	✓
	Public Wages and Salaries	0.127	✗	✓	✓	✗	✗	✓
Non-residential Fixed Investment	Industrial Production	0.042	✗	✗	✗	✓	✗	✓
	Nondefense Cap. Goods Shipments	0.018	✗	✗	✗	✓	✗	✓
	Oil Price	0.025	✗	✗	✗	✓	✗	✓
Residential Fixed Investment	All Employees, Res. Bldg. Const.	0.900	✗	✗	✗	✗	✓	✓
<i>In-sample Fit</i>								
N params.			60	62	69	72	64	85
BIC			-61877	-66990	-72087	-72044	-65328	-85687
AIC			-62165	-67288	-72419	-72390	-65635	-86096
<i>St. dev. of GDP estimate (percent)</i>								
All			0.271	0.268	0.266	0.270	0.272	0.266
2020s			0.176	0.174	0.173	0.175	0.185	0.181
2010s			0.151	0.148	0.147	0.151	0.152	0.148
2000s			0.159	0.154	0.153	0.159	0.160	0.154
1990s			0.201	0.197	0.195	0.201	0.200	0.194
1980s			0.244	0.241	0.239	0.244	0.245	0.240
1970s			0.278	0.274	0.272	0.278	0.278	0.272
1960s			0.319	0.314	0.312	0.319	0.319	0.312
1950s			0.597	0.597	0.595	0.592	0.597	0.589

Table 9: Alternative models

Table reports performance of alternate models. Upper half of the table lists specification details: extra variables considered, and the total number of Kalman parameters resulting. Lower half reports in-sample fit, as measured either by BIC or AIC, as well as the width of the confidence intervals. All models include all the observation variables listed in Table 2. pseudo- R^2 is that implicit in the bivariate regression in monthly terms, equivalent to that reported in Table 4. Sample: Jan 1959-Dec 2024

7 Conclusions

In this paper, I estimated monthly national accounts series consistent with the quarterly NIPA data, exactly monthly data for components, and other monthly indicators. I provide both point estimates and confidence intervals, and validate their coverage on known data. I also conduct simple analysis of recession dates and business cycle statistics using this data.

There may be ways to improve on the estimates provided here, either by using alternate statistical frameworks or by using extra data sources. However, the narrowness of the confidence intervals for the constructed series suggests that any potential gains are likely small.

In any case, the constructed series are precise enough to be useful in a wide range of applied research, opening up other avenues for future work where the lack of monthly national accounts data would otherwise be a constraint.

References

- Akbal, Omer Faruk, Mr Seung M Choi, Seung Choi, Mr Futoshi Narita, and Jiaxiong Yao**, *Panel Nowcasting for Countries Whose Quarterly GDPs are Unavailable*, International Monetary Fund, 2023.
- Andrews, Donald WK and Biao Lu**, “Consistent model and moment selection procedures for GMM estimation with application to dynamic panel data models,” *Journal of econometrics*, 2001, 101 (1), 123–164.
- Aruoba, S Boragan, Francis X Diebold, and Chiara Scotti**, “Real-time measurement of business conditions,” *Journal of Business & Economic Statistics*, 2009, 27 (4), 417–427.
- Beyer, Robert CM, Yingyao Hu, and Jiaxiong Yao**, *Measuring quarterly economic growth from outer space*, International Monetary Fund, 2022.
- Brave, Scott A, R Andrew Butters, David Kelley et al.**, “A new “big data” index of US economic activity,” *Economic Perspectives, Federal Reserve Bank of Chicago*, 2019, 1.
- Hu, Yingyao and Jiaxiong Yao**, “Illuminating economic growth,” *Journal of Econometrics*, 2022, 228 (2), 359–378.
- Koop, Gary, Stuart McIntyre, James Mitchell, and Aubrey Poon**, “Reconciled estimates of monthly GDP in the United States,” *Journal of Business & Economic Statistics*, 2023, 41 (2), 563–577.

Mariano, Roberto S and Yasutomo Murasawa, “A new coincident index of business cycles based on monthly and quarterly series,” *Journal of Applied Econometrics*, 2003, 18 (4), 427–443.

Stanger, Mr Michael, *A monthly indicator of economic growth for low income countries*, International Monetary Fund, 2020.

Stock, James H and Mark W Watson, “A probability model of the coincident economic indicators,” 1988.

– **and** –, “New indexes of coincident and leading economic indicators,” *NBER macroeconomics annual*, 1989, 4, 351–394.

– **and** –, “Distribution of quarterly values of GDP/GDI across months within the quarter,” *Research Memorandum*]. Retrieved from https://www.princeton.edu/~mwatson/mgdp_gdi/Monthly_GDP_GDI_Sept20.pdf, 2010.

A Baseline: Extra charts and tables

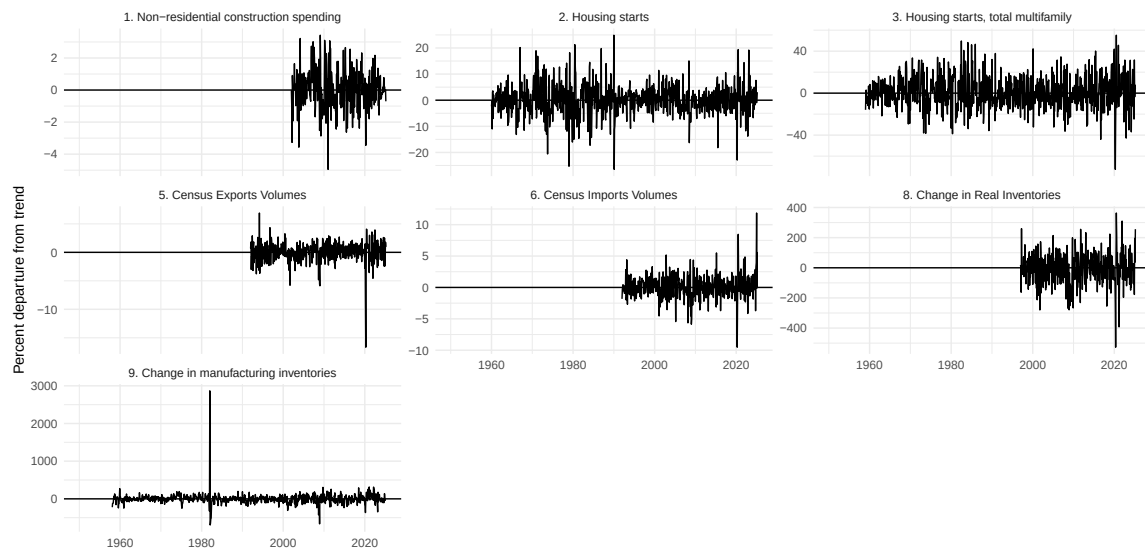


Figure A1: Monthly Indicators

Figure shows residualized monthly indicators.

	GDP	Services	Goods	NRFI	RFI	Inv.	Exports	Imports	Gov.
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)
Order 1 time polynomial	11.375*** (0.114)	12.292*** (0.114)	12.492*** (0.159)	17.056*** (0.169)	6.916*** (0.274)	178.464* (92.866)	20.733*** (0.446)	21.377*** (0.513)	7.340*** (0.174)
Order 2 time polynomial	-1.085*** (0.098)	-1.742*** (0.100)	-0.036 (0.135)	-0.718*** (0.198)	-1.941*** (0.392)	-4.049 (66.321)	-1.643*** (0.524)	-1.337** (0.557)	-1.178*** (0.174)
Constant	9.062*** (0.007)	8.256*** (0.007)	7.382*** (0.012)	6.579*** (0.012)	6.191*** (0.018)	33.880*** (5.339)	6.213*** (0.031)	6.402*** (0.036)	7.667*** (0.013)
Observations	301	301	301	301	301	301	301	301	301
R ²	0.998	0.999	0.995	0.993	0.840	0.037	0.991	0.990	0.977

Note:

*p<0.1; **p<0.05; ***p<0.01

Table A1: Estimated coefficients for the polynomial trends

Table reports the estimated trend coefficients from the quarterly data, the $\tilde{\gamma}^i$ of equation (14) for each of the GDP components. Sample: 1950Q1-2024Q4.

Variable	$\sqrt{Var\mu_t^i}$	$\sqrt{Var_t\mu_{t+1}^i}$	$\sqrt{\Sigma_t^{i,i}}$	Ratios	
	(1)	(2)	(3)	(1)/(3)	(2)/(3)
GDP, sum of components	3.81	1.23	0.26	14.5	4.7
Services consumption	2.59	0.52	0.03	75.1	15.0
Goods consumption	4.98	1.02	0.07	73.0	14.9
Non-residential Fixed Investment	8.07	2.22	0.86	9.3	2.6
Residential Fixed Investment	18.14	4.32	1.60	11.4	2.7
Government	6.45	1.90	0.89	7.2	2.1
Exports	11.46	2.65	0.97	11.8	2.7
Imports	12.52	2.30	0.72	17.5	3.2
Inventories	138.51	145.46	120.35	1.2	1.2

Table A2: Signal-to-noise ratios: GDP components, sample average

Table shows averages by decade of the standard deviation of the (1) unconditional and (2) one-step-ahead forecasts of monthly national accounts series, and well as the model-implied standard error (3). Units are in percentage of trend. Sample: Jan 1950-Mar 2025.

Decade	Services	Goods	NRFI	RFI	Gov.	Exports	Imports	Inv.
<i>Quarterly</i>								
1960s	0.88	0.95	0.91	0.15	0.68	-0.20	0.66	0.59
1970s	0.69	0.85	0.75	0.67	-0.03	0.32	0.65	0.55
1980s	0.73	0.76	0.39	0.69	0.53	0.43	0.79	0.59
1990s	0.87	0.94	0.92	0.81	-0.30	0.64	0.91	0.50
2000s	0.86	0.84	0.81	0.66	-0.59	0.64	0.94	0.79
2010s	0.76	0.58	0.47	0.24	0.15	-0.21	-0.02	0.42
2020s	0.94	0.52	0.81	0.29	-0.23	0.83	0.81	0.60
All	0.77	0.77	0.72	0.54	0.29	0.40	0.73	0.42
<i>Monthly</i>								
1960s	0.81	0.93	0.82	0.27	0.74	-0.31	0.63	0.57
1970s	0.64	0.88	0.58	0.73	0.23	0.08	0.75	0.46
1980s	0.80	0.84	0.23	0.83	0.66	0.24	0.83	0.50
1990s	0.84	0.92	0.83	0.75	-0.13	0.50	0.81	0.50
2000s	0.84	0.84	0.78	0.71	-0.55	0.58	0.94	0.78
2010s	0.79	0.57	0.42	0.19	0.23	-0.13	0.09	0.40
2020s	0.93	0.62	0.83	0.43	0.12	0.81	0.83	0.48
All	0.63	0.63	0.52	0.45	0.46	0.25	0.58	0.32

Table A3: Business-cycle statistics: Correlation with GDP

Table shows correlations of GDP with each component at business cycle frequencies. All series are first filtered using a Christiano-Fitzgerald filter to eliminate fluctuations at horizons longer than 5 years. “Services” and “Goods” columns are private consumption of services and goods respectively. Sample: Jan 1960-Mar 2025.

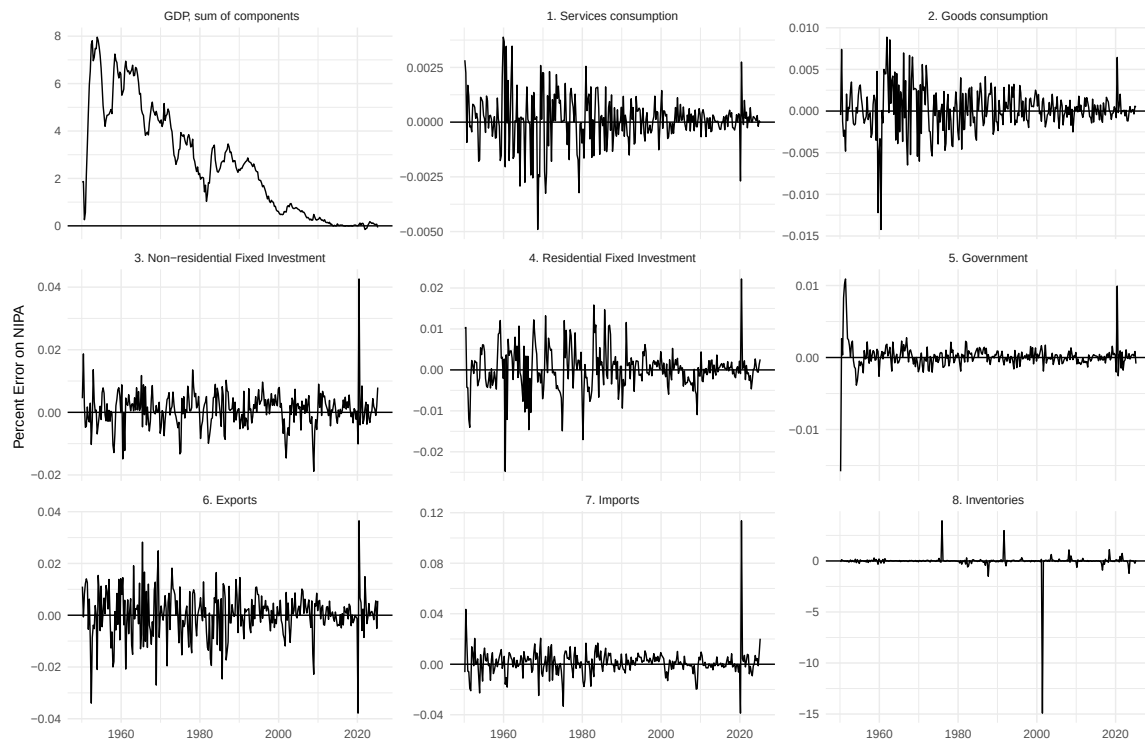


Figure A2: Error on sum of monthly series versus NIPA data: Baseline model.

Figure shows percentage error on GDP and components for the quarterly sums of the monthly series. Here “GDP” is the sum of components only, excluding the aggregation error. Sample: Jan 1950-Mar 2025.

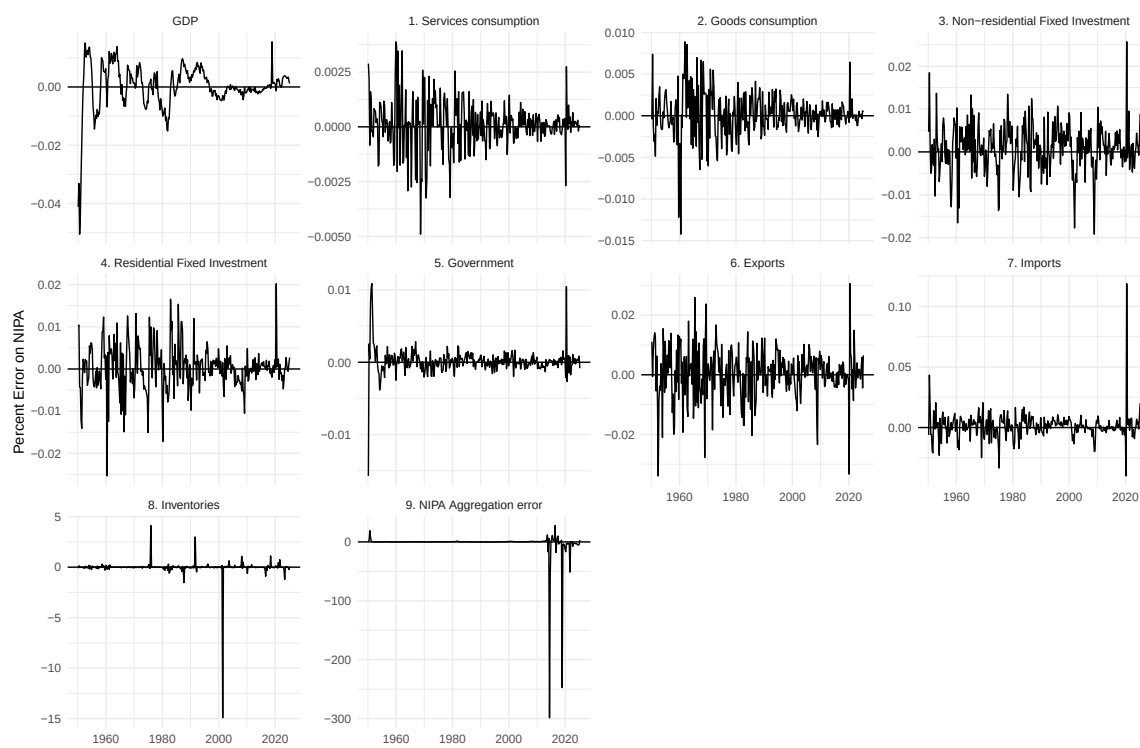


Figure A3: Error on sum of monthly series versus NIPA data: Model with aggregation error,

Figure shows percentage error on GDP and components for the quarterly sums of the monthly series. Here “GDP” includes the aggregation error. Sample: Jan 1950-Mar 2025.

B An algorithm for estimating monthly national account series

To summarize, the steps used to compute estimated monthly national accounts data are as follows.

1. Estimate trends and compute departures from trend as described in Section 4.1. This gives the trends $\bar{Y}_t, \bar{y}_t$ and departures from trend $\hat{Y}_t, \hat{y}_t$
2. Fit ARIMA processes to clean the predictors to form the $\epsilon_t^{i,j}$.
3. Estimate the parameters of the Kalman framework by GMM as described in Section 4.4. Denote the GMM point estimate by Ψ^* .
4. Compute the point estimate and Kalman prediction error covariance for $\hat{y}_t$ using a Kalman smoother with parameters Ψ^* . The point estimate for the components, μ_t , is complete.
5. Multiply the estimates in step 4 by their trends to obtain the point estimate and filter covariances for y_t .
6. Calculate the numerical derivative of the point estimate by perturbing each entry in Ψ^* by a small amount and repeating steps 4 and 5.
7. Multiply the derivative in step 6 by the asymptotic parameter covariance of Ψ^* from GMM estimation and add to the Kalman prediction error covariance computed in step 4. This is the final sequence of covariance estimates of the GDP components, Σ_t .
8. Compute the point estimate and variance of GDP using equation (3).

C Other episodes

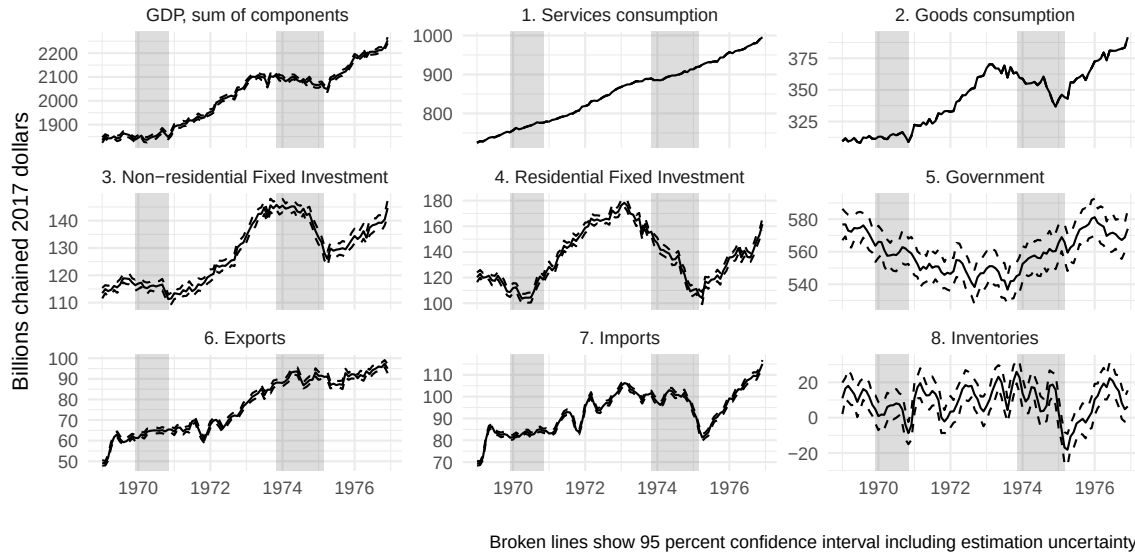


Figure A4: Estimated Monthly National Accounts Series: Early 1970s

Figure shows estimated monthly national accounts series plus 95 percent confidence interval. Shaded areas are NBER recessions. Sample: Jan 1969-Dec 1976.

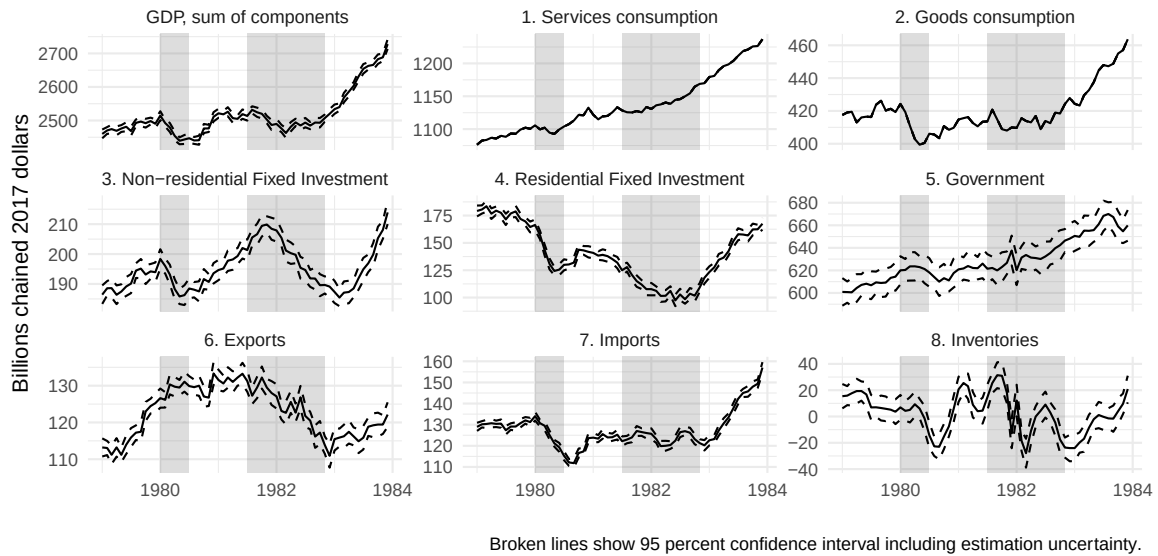


Figure A5: Estimated Monthly National Accounts Series: Volkler disinflation

Figure shows estimated monthly national accounts series plus 95 percent confidence interval. Shaded areas are NBER recessions. Sample: Jan 1979-Dec 1983.

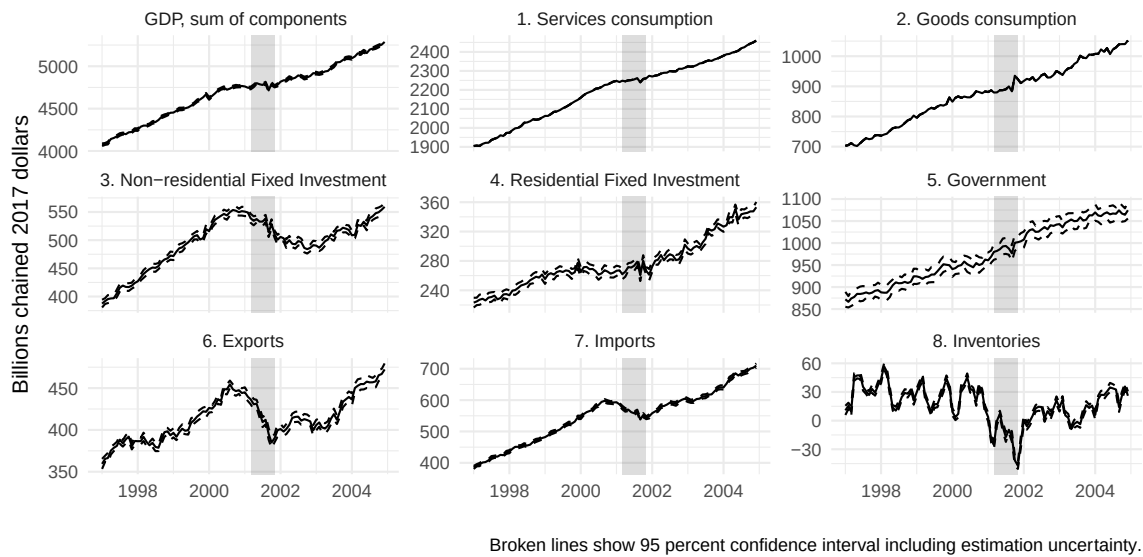


Figure A6: Estimated Monthly National Accounts Series: Late 1990s/early 2000s

Figure shows estimated monthly national accounts series plus 95 percent confidence interval. Shaded areas are NBER recessions. Sample: Jan 1997-Dec 2004.

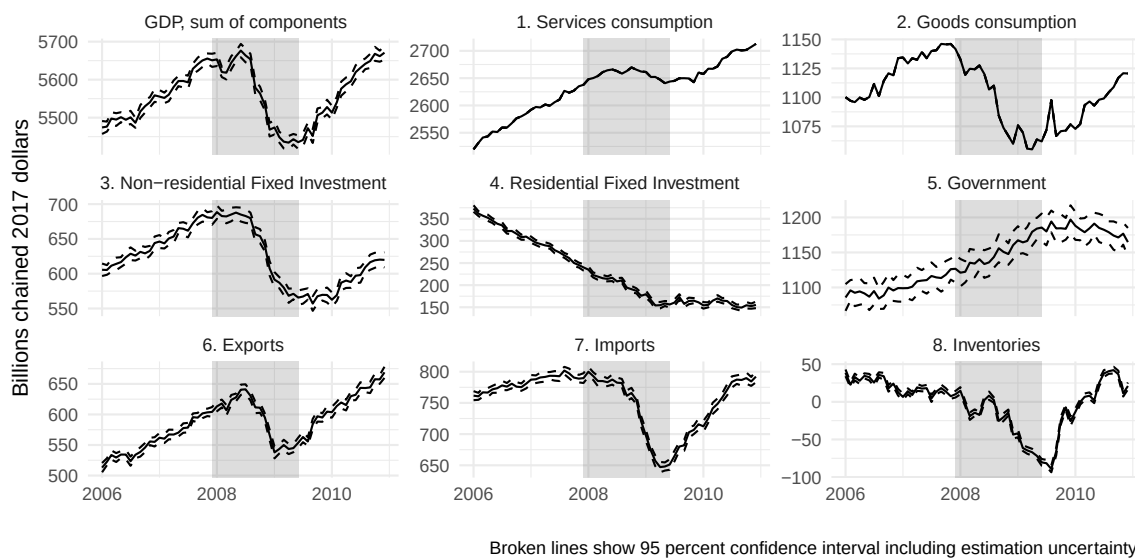


Figure A7: Estimated Monthly National Accounts Series: Global Financial Crisis

Figure shows estimated monthly national accounts series plus 95 percent confidence interval. Shaded areas are NBER recessions. Sample: Jan 2006-Dec 2010.

D Robustness

	Decade	GDP	Exports	Goods	Gov.	Imports	Inv.	NRFI	RFI	Services
Efficient GMM	1950s	-0.000	-0.000	-0.000	-0.000	0.000	0.388	-0.000	0.000	-0.000
	1960s	-0.000	-0.001	-0.000	-0.000	0.000	-2.358	-0.001	0.002	0.000
	1970s	-0.000	-0.001	0.000	-0.000	0.000	7.305	-0.001	0.001	-0.000
	1980s	-0.000	-0.002	0.000	-0.000	0.001	-24.164	-0.001	0.002	-0.000
	1990s	-0.000	-0.002	-0.000	-0.000	-0.000	-1.792	-0.001	-0.001	0.000
	2000s	-0.000	-0.001	0.000	0.000	0.000	-0.088	-0.001	0.001	-0.000
	2010s	-0.000	-0.001	-0.000	0.000	-0.000	-0.871	0.000	-0.000	0.000
	2020s	-0.001	0.007	0.000	-0.002	0.009	0.660	-0.012	-0.008	-0.000
Fitting 12 AR equations	1950s	-0.000	0.000	-0.000	0.000	-0.000	-2.706	-0.000	-0.000	-0.000
	1960s	0.000	0.001	-0.000	0.000	0.000	4.118	-0.000	0.002	-0.000
	1970s	0.000	0.001	-0.000	0.000	0.000	-11.276	-0.000	0.001	-0.000
	1980s	0.000	0.002	-0.000	0.000	0.000	1.970	-0.000	0.002	0.000
	1990s	-0.000	0.000	0.000	0.000	-0.000	-1.503	-0.000	0.001	-0.000
	2000s	-0.000	0.001	-0.000	0.000	0.000	0.156	-0.000	0.001	0.000
	2010s	0.000	0.000	-0.000	0.000	0.000	-0.342	-0.000	0.001	0.000
	2020s	0.000	0.009	-0.000	0.002	0.011	0.025	-0.002	0.009	-0.000
Third-order trend	1950s	0.000	0.002	0.000	0.000	-0.000	-2.195	0.000	0.000	0.000
	1960s	0.000	0.003	-0.000	-0.000	-0.002	7.398	-0.001	-0.001	0.000
	1970s	0.000	0.003	-0.000	-0.000	-0.001	-28.436	-0.000	0.000	0.000
	1980s	0.000	0.005	0.000	-0.000	-0.001	14.672	-0.001	-0.004	-0.000
	1990s	0.000	-0.001	0.000	0.000	-0.000	-6.773	-0.001	-0.001	-0.000
	2000s	0.000	-0.001	0.000	0.001	0.000	-1.929	-0.001	-0.003	-0.000
	2010s	-0.000	-0.000	0.000	0.000	0.000	-5.483	-0.000	0.000	-0.000
	2020s	-0.000	0.018	0.000	0.011	0.104	9.185	-0.001	0.017	-0.000
Including aggregation error	1950s	-5.593	0.000	-0.000	-0.002	-0.001	0.775	0.000	-0.000	0.000
	1960s	-4.747	0.001	0.000	0.000	-0.000	-17.491	-0.000	-0.002	0.000
	1970s	-3.372	0.001	0.000	0.000	0.000	-1.444	-0.001	-0.001	0.000
	1980s	-2.411	0.005	-0.000	0.000	-0.000	-103.083	-0.002	-0.003	-0.000
	1990s	-1.652	0.001	0.000	-0.001	-0.000	-6.316	-0.000	-0.001	0.000
	2000s	-0.801	-0.000	-0.000	0.000	0.002	0.756	-0.008	-0.000	-0.000
	2010s	-0.177	0.002	-0.000	0.003	0.010	-65.677	-0.044	0.000	-0.000
	2020s	0.121	0.003	-0.000	0.011	0.015	-6.907	0.015	0.014	-0.000

Table A4: Robustness to alternate specifications: Average difference versus baseline

Table shows average difference in monthly national accounts series compared to the baseline estimates for different model specifications. Differences are expressed in percent. Sample: Jan 1959-Mar 2025

	Decade	GDP	Exports	Goods	Gov.	Imports	Inv.	NRFI	RFI	Services
Efficient GMM	1950s	0.034	0.036	0.017	0.030	0.050	5.804	0.051	0.079	0.007
	1960s	0.042	0.112	0.000	0.085	0.070	4.974	0.152	0.153	0.000
	1970s	0.034	0.098	0.000	0.091	0.061	11.611	0.132	0.192	0.000
	1980s	0.042	0.146	0.000	0.093	0.089	28.685	0.179	0.184	0.000
	1990s	0.037	0.314	0.000	0.105	0.112	18.207	0.197	0.170	0.000
	2000s	0.039	0.298	0.000	0.082	0.062	2.628	0.282	0.175	0.000
	2010s	0.036	0.255	0.000	0.087	0.070	2.602	0.208	0.194	0.000
Fitting 12 AR equations	2020s	0.102	0.543	0.000	0.158	0.163	2.208	0.438	0.237	0.000
	1950s	0.033	0.031	0.021	0.009	0.056	7.498	0.050	0.082	0.009
	1960s	0.011	0.091	0.000	0.019	0.025	5.462	0.021	0.066	0.000
	1970s	0.010	0.084	0.000	0.019	0.029	13.918	0.016	0.056	0.000
	1980s	0.009	0.121	0.000	0.022	0.036	3.172	0.015	0.071	0.000
	1990s	0.012	0.114	0.000	0.017	0.034	5.222	0.061	0.044	0.000
	2000s	0.010	0.120	0.000	0.031	0.018	0.443	0.053	0.051	0.000
Third-order trend	2010s	0.010	0.077	0.000	0.019	0.017	0.785	0.041	0.041	0.000
	2020s	0.028	0.258	0.000	0.070	0.095	0.896	0.089	0.101	0.000
	1950s	0.054	0.113	0.030	0.033	0.132	14.362	0.056	0.126	0.019
	1960s	0.035	0.242	0.000	0.066	0.408	12.489	0.114	0.240	0.000
	1970s	0.029	0.222	0.000	0.072	0.338	35.160	0.095	0.261	0.000
	1980s	0.034	0.284	0.000	0.093	0.515	21.882	0.142	0.368	0.000
	1990s	0.067	0.300	0.000	0.276	0.642	26.160	0.113	0.289	0.000
Including aggregation error	2000s	0.034	0.269	0.000	0.449	0.629	9.825	0.160	0.362	0.000
	2010s	0.029	0.228	0.000	0.394	0.648	9.674	0.108	0.486	0.000
	2020s	0.054	0.524	0.000	0.644	1.276	21.307	0.162	0.454	0.000
	1950s	13.304	0.296	0.182	1.409	1.017	31.257	1.133	0.367	0.282
	1960s	7.730	0.480	0.000	0.884	0.222	32.383	1.225	0.104	0.000
	1970s	4.742	0.376	0.000	0.719	0.157	30.967	0.895	0.088	0.000
	1980s	5.251	0.758	0.000	1.331	0.332	174.000	1.719	0.178	0.000
	1990s	4.419	0.539	0.000	1.912	0.666	53.240	2.574	0.185	0.000
	2000s	2.700	0.415	0.000	2.550	1.071	7.617	2.854	0.112	0.000
	2010s	2.488	2.483	0.000	17.431	8.373	95.004	18.141	0.381	0.000
	2020s	1.864	1.373	0.000	7.900	4.448	30.344	8.800	0.297	0.000

Table A5: Robustness to alternate specifications: Average absolute difference versus baseline

Table shows average absolute difference in monthly national accounts series compared to the baseline estimates for different model specifications. Differences are expressed in percent. Sample: Jan 1959-Mar 2025



PUBLICATIONS

Estimated Monthly National Accounts for the United States
Working Paper No. WP/2025/134