

UNDERSTANDING UNCERTAINTY SHOCKS

Anna Orlik and Laura Veldkamp *

December 14, 2012

Abstract

For decades, macroeconomists have searched for shocks that are plausible drivers of business cycles. A recent advance in this quest has been to explore uncertainty shocks. Researchers use a variety of forecast and volatility data to justify heteroskedastic shocks in a model, which can then generate realistic cyclical fluctuations. But the relevant measure of uncertainty in most models is the conditional variance of a forecast. When agents form such forecasts with state, parameter and model uncertainty, neither forecast dispersion nor innovation volatilities are good proxies for conditional forecast variance. We use observable data to select and estimate a forecasting model and then ask the model to inform us about what uncertainty shocks look like and why they arise.

Some times feel like uncertain times for the aggregate economy. At other times, events appear to be predictable, volatility is low, confidence is high. An active emerging literature argues that changes in uncertainty can explain business cycle fluctuations, stock prices, and banking crises. Uncertainty shocks are typically modeled as volatility shocks. Authors assume that a shock (e.g. productivity) is drawn from a distribution whose variance changes over time. Sometimes the high and low volatility are calibrated to match the size of business cycles. Other times, they are matched to features of macro forecast data, the volatility index (VIX) or uncertainty proxies.

This paper seeks to clarify the relationship between uncertainty and commonly-used measures of uncertainty. We show that when agents consider the possibility that their model might be wrong, uncertainty, volatility and forecast errors are conceptually distinct.

*Preliminary draft. Please send comments to anna.a.orlik@frb.gov and lveldkam@stern.nyu.edu. Many thanks to Isaac Baley, David Johnson and Nic Kozeniauskas for outstanding research assistance. We acknowledge grant assistance from the Stern Center for Global Economy and Business. The views expressed herein are those of the authors and do not necessarily reflect the position of the Board of Governors of the Federal Reserve or the Federal Reserve System.

Estimating this forecasting model on GDP data allows us to infer a process for uncertainty. By comparing the uncertainty process to volatility and forecast errors, we learn that model and parameter uncertainty play a large role in uncertainty shocks.

We define macroeconomic uncertainty as the variance of next-period GDP growth y_{t+1} , conditional on all information observed through time- t : $Var[y_{t+1}|\mathcal{I}_t]$. We use this definition because in most models, this is the theoretically-relevant moment. When there is an option value of waiting, beliefs with a higher conditional variance (imprecise beliefs) raise the value of waiting to observe additional information. Thus, it is uncertainty in the form of a higher conditional variance that typically delays consumption or investment and thus depresses economic activity. Since this conditional variance is not directly observable in data, our objective is to infer this uncertainty with a forecasting model, and compare it to commonly-used proxy variables to assess whether any of those proxies accurately capture the timing and magnitude of uncertainty shocks.

A key premise of our paper is that we should use a model to infer what uncertainty is. Why not just estimate a stochastic volatility, for example, and then call the variance of the innovations uncertainty? Indeed, this would be uncertainty for an agent who knew that his model was the true GDP process and knew its true parameters. But the macroeconomy is not governed by a simple, known model and we surely don't know its parameters. Instead, we estimate simple models to approximate complex processes and constantly use new data to update the beliefs. In such a setting, uncertainty and volatility can behave quite differently. Two examples illustrate this difference. In the first example, uncertainty can change with constant volatility. When a shock hits that is low probability under the agent's current model and parameters, it will make the agent shift probability weight to other models and/or parameters. Being less certain about the data-generating process raises uncertainty about future GDP growth. In the second example, a surge in volatility can have a tiny effect on uncertainty. Suppose that our forecasting agent receives a signal, s_t , that is future gdp growth y plus noise $s_t = y_{t+1} + \epsilon_t$. If the variance of noise $var(\epsilon_t)$ is constant, then more volatile GDP (high $var(y_{t+1})$) means the signal-to-noise ratio of s_t is higher. Because GDP fluctuates more than the noise ϵ , it is easier to discern. If signal precision is high, relative to the precision of prior beliefs about GDP growth, then even a large volatility shock could have a very small effect on uncertainty because it raises signal informativeness

at the same time. Similarly, forecast errors and dispersion are related to uncertainty, but are not equivalent. This paper builds and estimates models that capture these kinds of effects. By comparing model-generated measures of uncertainty to estimates of volatility, we can learn about just how different volatility, forecast dispersion and uncertainty are.

Inferring uncertainty from a model also provides insight about where uncertainty shocks come from. The common practice is to model uncertainty shocks like an aggregate belief shock, where all agents wake up one day knowing that the world will be a more volatile place. This is similar to assuming a common preference shock. Without any discipline on beliefs, a model can explain almost anything. In this paper, uncertainty is the outcome of a Bayesian belief formation process, disciplined by a model estimated on 45-60 years of GDP data. Through the model, we can understand what kinds of events or sequences of events cause uncertainty to be high and why.

Since our estimates of uncertainty will be model-dependent, we examine a set of simple forecasting models that capture some key features of professional forecasts. To build up intuition for how the model works, we start with one of the simplest settings: a 2-state regime switching model of GDP growth. Then, we start adding features to the model to make it more realistic. Each period t , our forecaster observes time- t GDP growth and uses the complete history of GDP data to estimate the state and forecast the GDP growth in $t + 1$. Since forecasters are not endowed with knowledge of model parameters, we allow them to re-estimate all the parameters of the model each period. This model generates small fluctuations in uncertainty, in part because we assumed a normal distribution of GDP growth shocks. Normal variables have the property that the conditional variance does not depend on the conditional mean. In order to give our model a better chance of generating uncertainty shocks, we allow the forecaster to entertain the possibility that the distribution of shocks to GDP growth has fat tails. Thus, the next variant of the model allows for model uncertainty: One model has normal shocks, the other model has t-distributed shocks and the agent estimates the probability of each model being the true data-generating process. Finally, we notice that with only two states, the GDP estimates have higher average errors than the forecasters do. One way to bring the model and data closer on this dimension is to give forecasters additional signals about future GDP. This allows the model to speak to forecast dispersion and to achieve a closer match with the moments of the forecaster data.

Our model also teaches us how uncertainty shocks can arise when rational, Bayesian agents use simple models to forecast complex economic processes. Each of these models generates a time-series of uncertainty. At each date t , the forecaster forms his forecast of y_{t+1} and computes the variance of that estimate $var[y_{t+1}|y^t]$. By examining this series, we gain insight about what causes uncertainty to rise. When agents use simple models to forecast a complex process, there will be times when the model performs poorly, inducing agents to question the model and its parameters. The model uncertainty, combined with parameter learning, cause a sudden spike in forecast uncertainty. This seems to happen primarily in recessions, because recessions are a time when there is a series of quarterly growth outliers. But uncertainty also has a very slow-moving, persistent component that comes from the fact that uncertainty is only resolved slowly. It takes numerous, informative observations for uncertainty to resolve.

We compare our model-based uncertainty series to commonly-used uncertainty proxies and find that it is less variable, but more persistent than the proxy variables. The best proxies are either forecast errors or forecast dispersion. But neither achieves a high degree of predictive power.

Most of the paper considers uncertainty about GDP because there are GDP forecasts that can be used to evaluate model performance. But most uncertainty-shock theories have GDP as an endogenous outcome and instead use productivity (TFP) as the driving process. Uncertainty shocks are changes in TFP uncertainty. Therefore, the last section of the paper (section 4) uses the same tools and insights that we have built up to infer GDP uncertainty to create a time series process for uncertainty about productivity. This uncertainty process can then be used to fit an uncertainty-driven business cycle model.

Related literature A new and growing literature uses uncertainty shocks as a driving process to explain business cycles (e.g., Bloom, Floetotto, Jaimovich, Sapora-Eksten, and Terry (2012), Basu and Bundick (2012) Christiano, Motto, and Rostagno (2012), Bianchi, Ilut, and Schneider (2012)), to explain investment dynamics Bachmann and Bayer (2012a), to explain asset prices (e.g., Bansal and Shaliastovich (2010), Pastor and Veronesi (2012)), and to explain banking panics (Bruno and Shin, 2012). These papers are complementary to ours. We explain where uncertainty shocks come from, while these papers trace of the

economic and financial consequences of the shocks.¹

More similar to our exercise is a set of papers that measure uncertainty shocks in various ways. Bloom (2009), Baker, Bloom, and Davis (2012), Fernandez-Villaverde, Kuester, Guerron, and Rubio-Ramirez (2012), Fernandez-Villaverde, Guerron-Quintana, Rubio-Ramirez, and Uribe (2011), Nakamura, Sergeyev, and Steinsson (2012) and Born, Peter, and Pfeifer (2011) document the properties of uncertainty shocks in the U.S. and in emerging economies, while Bachmann, Elstner, and Sims (2012) use forecaster data to measure ex-ante and ex-post uncertainty in Germany. While our paper also engages in a measurement exercise, what we add to this literature is a quantitative model of how and why such shocks arise.

The theoretical part of our paper grows out of an existing literature that estimates Bayesian forecasting models with model uncertainty. Cogley and Sargent (2005) use such a model to understand the behavior of monetary policy, while Johannes, Lochstoer, and Mou (2011) estimate a similar type of model on consumption data to capture properties of asset prices. While the mechanics of the model estimation is similar, the focus on uncertainty shocks and the finding that such a model can quantitatively explain the extent of uncertainty shocks in survey data distinguish our paper. Nimark (2012) also generates increases in uncertainty by assuming that only outlier events are reported. Thus, the publication of a signal conveys both the signal content and information that the true event is far away from the mean. Such signals can increase agents' uncertainty. But that paper does not attempt to quantitatively explain the fluctuations in uncertainty measures. Our paper is in the spirit of Hansen's AEA presidential address (Hansen, 2007) which advocates putting agents in the model on equal footing with an econometrician who is learning about his environment over time and struggling with model selection.

1 A Forecasting Model

The purpose of the model is to explain how much and why uncertainty varies. In the model, an agent will observe a time-series of data and forecast the next period's realization of that time-series. Uncertainty in our forecasting model is defined as the expected squared

¹In contrast, Bachmann and Bayer (2012b) argue that there is little impact of uncertainty on economic activity.

difference between the forecast and the realization. In other words, it is the conditional variance of the forecast.

We model an agent who acts like an econometrician, in the spirit of Hansen (2007). He observes a series of data and uses Bayesian updating to estimate the state (probability of being in recession), the parameters, and the probability that innovations are normally or t-distributed. For a problem with this many interdependent objects to learn about, inference is tricky. We need to use a Metropolis-Hastings algorithm to estimate the model objects. Those estimates allow the agent to forecast next period’s productivity. His expected forecast error determines his degree of uncertainty. The next period, the agent sees a new piece of data, re-estimates states, parameters and the model probabilities, and as a result, is faced with a new degree of uncertainty. This process generates forecast errors that look similar to the errors of professional forecasters in the data. But the uncertainty estimates look quite different from forecast errors.

In order to fully understand what elements of the model are responsible for which features of the uncertainty process, we build up the model in three stages. In the first stage, we endow the agent with knowledge of the model and the model parameters and only ask him to estimate the probability of being in one of two states and forecast the probability of being in each state tomorrow.

1.1 Data Description

There are two pieces of data that we use to evaluate and estimate our forecasting model. The first is GDP data from the Bureau of Economic Analysis. The variable we denote y_t is the growth rate of GDP. Specifically, it is the log-difference of the seasonally-adjusted real GDP series, times 400, so that it can be interpreted as an annualized percentage change.

We use the second set of data, professional GDP forecasts, to evaluate our forecasting models. We describe below the four key moments that we use to make that assessment. The data come from the Survey of Professional Forecasters, released by the Philadelphia Federal Reserve. The data are a panel of individual forecaster predictions of real US output for both the current quarter and for one quarter ahead from quarterly surveys from 1968 Q4 to 2011 Q4. In each quarter, the number of forecasters varies from quarter-to-quarter, with an average of 40.5 forecasts per quarter.

Formally, $t \in \{1, 2, \dots, T\}$ is the quarter in which the survey of professional forecasters is given. Let $i \in \{1, 2, \dots, I\}$ index a forecaster and $I_t \subset \{1, 2, \dots, I\}$ be the subset of forecasters who participate in a given quarter. Thus, the number of forecasts made at time t is $N_t = \sum_{i=1}^I \mathbb{I}(i \in I_t)$. Finally, let y_{t+1} denote the GDP growth rate over the course of period t . Thus, if GDP_t is the GDP at the end of period t , observed at the start of quarter $t + 1$, then $y_{t+1} \equiv \ln(GDP_t) - \ln(GDP_{t-1})$. This timing convention may appear odd. But we date the growth $t + 1$ because it is not known until the start of date $t + 1$. The growth forecast is constructed as $E_{it}[y_{t+1}] = \ln(E_{it}[GDP_t]) - \ln(GDP_{t-1})$.

A forecast of period- t GDP growth made at the start of period t , by forecaster i is denoted $E_{it}[y_t]$. Next, we define the average forecast error, a quantity that we will use to help select a forecasting model that performs well.

Definition 1. *A forecast error is the distance between a forecast of the growth rate and the realized growth rate:*

$$FE_{t+1} = \frac{\sum_{i \in I_t} |E_{it}[y_{t+1}] - y_{t+1}|}{N_t}. \quad (1)$$

1.2 A General Forecasting Model

We will examine a variety of models, each denoted $\mathcal{M}$. All share the following common structure. Let $\{y_i\}_{i=0}^t$ denote a time series data available to the forecaster at time t . We postulate the following two-state hidden Markov regime switching process drives the dynamics of y_t

$$y_t = \mu_{S_t} + \sigma e_t \quad (2)$$

where $e_{t+1} \sim \phi$. Models will differ in their probability distribution ϕ and in the information set available to forecasters. In every model, the information set will include y^t , the history of y observations up to and including time t . The state $S_t \in \{1, 2\}$ is a two-state Markov variable. State transitions are governed by an 2×2 transition probability matrix whose elements are $q_{ij} \equiv \Pr(S_t = j | S_{t-1} = i)$ with $\sum_j q_{ij} = 1 \ \forall i$. The transition matrix controls the persistence and stochastic evolution of the mean of the y_t process.

Each model has a vector θ of parameters. The elements of θ are the variance parameter σ , two state-dependent means μ_1, μ_2 , and the transition probabilities $q_{11}, q_{12}, q_{21}, q_{22}$.

The state S_t is never observed. Thus (2) is the observation equation of a filtering problem. It maps an unobservable state S_t into an observable outcome y_t . The standard deviation of the unexpected innovations to this process is what we call volatility. Note that when we compute what is an unexpected innovation, we take the model and its parameters as given.

Definition 2. *Volatility is the standard deviation of the unexpected innovations to the observation equation, taking the model and its parameters as given:*

$$VOL_t = \sqrt{E \left[(y_{t+1} - E[y_{t+1}|y^t, \theta, \mathcal{M}])^2 \middle| y^t, \theta, \mathcal{M} \right]}.$$

The agent, who we call a forecaster and index by i , is not faced with any economic choices. He simply uses Bayes' law to forecast future y outcomes. Specifically, at each date t , the agent conditions on his information set $\mathcal{I}_{it}$ and forms beliefs about y_{t+1} . We call the expected value $E(y_{t+1}|\mathcal{I}_{it})$ and agent i 's *forecast* and the square root of the conditional variance $Var(y_{t+1}|\mathcal{I}_{it})$ is what we call *uncertainty*.

Definition 3. *Uncertainty is the standard deviation of the time-($t + 1$) GDP growth, conditional on an agent's time- t information: $U_{it} = \sqrt{E \left[(y_{t+1} - E[y_{t+1}|\mathcal{I}_{it}])^2 \middle| \mathcal{I}_{it} \right]}$.*

Forecasters' forecasts will differ from the realized growth rate. This difference is what we call a forecast error.

Definition 4. *An agent i 's forecast error is the distance, in absolute value, between the forecast and the realized growth rate: $FE_{i,t+1} = |y_{t+1} - E[y_{t+1}|\mathcal{I}_{it}]|$.*

We date the forecast error $t + 1$ because it depends on a variable y_{t+1} that is not observed at time t . This is an ex-post measure because it is not measurable at the time when the forecast is made. Volatility and uncertainty are both ex-ante measures because they are time- t expectations of $t + 1$ outcomes, which are time- t measurable.

It is true that $U_{it} = \sqrt{E_t[FE_{i,t+1}^2|\mathcal{I}_t]}$. So, the expected squared forecast error is the same as uncertainty. Of course, what people measure with forecast errors is not the *expected* squared forecast error. It is an average of realized squared forecast errors: $\sqrt{1/N \sum_i FE_{i,t+1}^2}$. If all agents in the model make the same forecast and have the same forecast error, $\sqrt{1/N \sum_i FE_{i,t+1}^2} = FE_{it}, \forall i$.

1.3 Forecasting with Only State Uncertainty (Model 1)

Many papers equate volatility, uncertainty and squared forecast errors. Thus, we begin with a setting where volatility and uncertainty are equivalent and both are equal to the square root of expected squared forecast errors. In this section only, we assume that the forecaster does know the model $\mathcal{M}$ and the parameters of that model θ .

Model 1 assumptions All agents have an identical information set: $\mathcal{I}_{it} = \{y^t, \theta, \mathcal{M}\}$, $\forall i$. Furthermore, the transitory innovations to the Markov process (2) are standard, normal variables: $e_{t+1} \sim N(0, 1)$.

By examining definitions 2 and 3, we can see that when $\mathcal{I}_{it} = \{y^t, \theta, \mathcal{M}\}$, uncertainty and volatility are equivalent. Thus, when a forecaster knows his model and its parameters with certainty, volatility measures uncertainty. But even in this model, neither is always equal to the forecast error since FE_{it} depends on a random $t + 1$ outcome.

To compute the process for uncertainty, we use Bayesian updating to form the conditional variance in definition 3, as follows: Since the history of states $\{S_1, \dots, S_t\}$ is not observed, the agent will use the Bayes rule to filter the hidden state. Each period, he will use the new y_t data to update his belief about which state might have been realized, in the following way. The agent is interested in computing the predictive data density $p(y_{t+1}|y^t)$ and using it to make following forecast

$$E(y_{t+1}|y^t) = \int y_{t+1} p(y_{t+1}|y^t) dy_{t+1}. \quad (3)$$

Let $\psi_{n,t} = \Pr(S_t = n|y^t)$ with $\sum_{n=1}^N \psi_{n,t} = 1$ be the belief about state n that the forecaster holds in a given period t . Then, the forecast is constructed as

$$\begin{aligned} E(y_{t+1}|y^t) &= \sum_{n=1}^2 \Pr(S_{t+1} = n|y^t) \int y_{t+1} p(y_{t+1}|y^t, S_{t+1} = i, S_t = j) dy_{t+1} \\ &= [q_{11}\psi_{1,t} + q_{21}(1 - \psi_{1,t})] \mu_1 + [q_{12}\psi_{1,t} + q_{22}(1 - \psi_{1,t})] \mu_2 \end{aligned} \quad (4)$$

where the last equality follows from the fact, that the distribution conditional on the current and future Markov state is given by $p(y_{t+1}|y^t, S_{t+1} = i, S_t = j) = \frac{1}{\sigma\sqrt{2\pi}} \exp\left[-\frac{(y_{t+1} - \mu_i)^2}{2\sigma^2}\right]$.

The conditional variance is then given by

$$Var(y_{t+1}|y^t) = [q_{11}\psi_{1,t} + q_{21}(1 - \psi_{1,t})] \{1 - [q_{11}\psi_{1,t} + q_{21}(1 - \psi_{1,t})]\} (\mu_1 - \mu_2)^2 + \sigma^2$$

Notice that even though the shock e_t is i.i.d normal and the distribution of y_t conditional on a Markov state and history of observations is normal, the distribution of future values y_{t+1} is not i.i.d., due to the time variation in the moments generated by the shifting Markov state variable. The variation in the posterior state belief $\psi_{i,t+1}$ is the key source of variation in uncertainty.

To complete the description of this forecasting problem with state uncertainty only, we need to determine the evolution of the state belief over time. The agent starts with a prior belief $\psi_{n,0}$ and updates it using the Bayes law:

$$\psi_{i,t+1} = \frac{\sum_{j=1}^N p(y_{t+1}|y^t, S_{t+1} = i, S_t = j) q_{ji} \psi_{j,t}}{\sum_{i=1}^N \sum_{j=1}^N p(y_{t+1}|y^t, S_{t+1} = i, S_t = j) q_{ji} \psi_{j,t}} \quad (5)$$

Comparing model 1 to forecast data

If this model with state uncertainty only produces forecasts that closely resemble those in the forecast data, then that would support the idea that volatility and uncertainty are approximately interchangeable.

The parameters that are in the forecaster's information set come from estimating the model with maximum-likelihood, on the full sample of data. The data is GDP growth rates 1968:Q4-2011:Q4. The procedure for computing forecasts and uncertainty is as follows: We endow the agent with knowledge of the parameters in table 1 and with the initial belief that there is an equal probability of being in each state ($\psi_{n,0} = 0.5$). Each period, the agent updates the current state belief using (5) and forecasts y_{t+1} using (4).

In the 2-state model, the agent believes that he is in the good (high-growth) state, most of the time. When low GDP growth is realized, the agent starts to place more probability weight on the bad, low-growth state. Being less certain of which state currently prevails makes the forecaster's uncertainty rise. Because of the negatively skewed pattern of GDP

Table 1: **Maximum Likelihood Parameters**

These are maximum likelihood estimates of (2), using real GDP growth data from 1968-2011.

Parameter	Value
μ_1	-2.47%
μ_2	3.55%
σ^2	$(2.9\%)^2$
q_{11}	0.69
q_{22}	0.96

growth, the agent puts high probability weight on the good state, which is frequently realized. This makes low-growth realizations more likely to generate increases in uncertainty. Thus a 2-state model can generate the pattern of counter-cyclical uncertainty observed in the data.

The problem with this model is that with only 2 states, the forecaster only attached probability weights to 2 possible growth estimates. The model doesn't allow him to adjust the states and attain more nuanced forecasts of future output. In the results in table 2, this shows up as a forecast error that is too large, too variable and too autocorrelated. Moreover, the forecasts don't track GDP growth very well. Their correlation is only 0.12, compared to 0.72 in the forecaster data. Furthermore, the agent's average GDP forecast is too low, causing the average forecast error to be much higher than in the data. So while the 2-state model provides simple laboratory for exploring how forecaster's update their beliefs in simple forecasting models, it is not the best framework for explaining how economic agents (professional forecasters) form expectations.

	data	state unc	param unc	signals	model unc
Avg forecast error	1.91%	2.53%	2.43%	1.96%	2.43%
Std forecast error	1.52%	2.48%	2.41%	1.40%	2.39%
Autocorr forecast error	0.15	0.28	0.17	0.13	0.21
Corr(forecast $_{t+1 t}$, GDP_{t+1})	0.72	0.12	0.25	0.74	0.25

Table 2: **Properties of forecasts in model and data.** Column 1 labeled 'state unc' uses equation (4) to forecast y_{t+1} . Columns 2, 3, 4 and 5 use equations (4), (7), (10) and (12) to forecast y_{t+1} . For the data and model 3 with heterogeneous signals, the standard deviation is the time-series standard deviation of the cross-section average, i.e. $std(1/N \sum_i |E_{it}[y_{t+1}] - y_{t+1}|)$. Similarly, the autocorrelation and correlation with GDP use the average forecast error or average forecast at each date t .

1.4 Forecasting with State and Parameter Uncertainty (Model 2)

So far, we have assumed that the forecaster does not observe economic regimes but that he knows all the parameter values. That is obviously unrealistic because we've used data from the entire sample period to estimate these parameters. A forecaster at any date before the last date in our sample would not have access to such information. Relaxing this assumption is important because if we want to understand uncertainty, part of uncertainty is not just whether there is a boom or recession, but also what the model parameters are. Solving a model where uncertainty about parameter values is a component of uncertainty helps us understand what role parameter uncertainty plays in producing uncertainty shocks. Therefore, the second forecasting model we consider is one where the agent forecasts using the same 2-state hidden Markov process (2), but the agent does not know any of the parameters θ and updates beliefs about them using Bayes' law.

Model 2 assumptions Each forecaster has an identical information set: $\mathcal{I}_{it} = \{y^t, \mathcal{M}\}, \forall i$. The model $\mathcal{M}$ is (2) with transitory innovations that are normally distributed: $e_{t+1} \sim N(0, 1)$.

By examining definitions 2 and 3, we can see that when $\mathcal{I}_{it} \neq \{y^t, \theta, \mathcal{M}\}$, uncertainty and volatility differ. Neither is equal to the realized forecast error, FE_t . The difference between the model 2 uncertainty process and the model 1 uncertainty process points to the importance of parameter uncertainty.

Our forecaster starts with a prior distribution of parameters, $p(\theta)$, described in table 10. Each period, the agent observes y_t , and updates his beliefs about parameters and states using Bayes' law. The posterior distribution $p(\theta, S^t | y^t)$ summarizes beliefs after observing the data $y^t = (y_1, y_2, \dots, y_t)$. Using Bayes' rule, we can rewrite $p(\theta, S^t | y^t) = p(\theta | y^t) p(S^t | y^t, \theta)$. Conditional on the history y^t and parameters θ , the agent can update, as in model 1 (equation 5). The challenge is to determine the posterior distribution of parameters $p(\theta | y^t)$. It is a high-dimensional object whose dependence on the data is not necessarily linear.

To compute posterior beliefs about parameters, we employ a Markov Chain Monte Carlo (MCMC) technique.² At each date t , the MCMC algorithm produces a sample of

²More details are presented in the Appendix. Also, see Johannes, Lochstoer, and Mou (2011) for a recursive implementation of a similar problem of sampling from the sequence of distributions.

possible parameter vectors, $\{\theta^d\}_{d=1}^D$, such that the probability of any parameter vector θ^d being in the sample is equal to the posterior probability of those parameters, $p(\theta|y^t)$. Therefore, we can compute an approximation to any integral by adding over sample draws: $\int f(\theta)p(\theta|y^t)d\theta \approx 1/D \sum_d f(\theta^d)$. Every parameter draw θ^d implies a probability of being in state i , denoted $\psi_{i,t}^d = \Pr(S_t = i|y^t, \theta^d)$. Thus, the forecaster can construct his forecast as

$$\begin{aligned} E(y_{t+1}|y^t) &= \sum_{S^{t+1}} \int y_{t+1} p(y_{t+1}|\theta, S^{t+1}) p(\theta, S^{t+1}|y^t) d\theta \\ &\approx \frac{1}{D} \sum_{d=1}^D \left[q_{11}^d \psi_{1,t}^d + q_{21}^d (1 - \psi_{1,t}^d) \right] \mu_1^d + \left[q_{12}^d \psi_{1,t}^d + q_{22}^d (1 - \psi_{1,t}^d) \right] \mu_2^d. \end{aligned} \quad (6)$$

$$(7)$$

Comparing model 2 to forecast data Including parameter uncertainty is not only more realistic. It also helps the forecasting model to come closer to the data. Column 3 of table 2 shows that adding parameter uncertainty reduces the average forecast error and its volatility. Although, the fit is improved, both moments are still far higher than what the data suggest. Parameter uncertainty did bring the autocorrelation of forecast errors in line with the data and it resulted in a doubling of the correlation of the forecast with GDP. Although, that correlation is still only one-third of what it is in the data.

1.5 A Model with Heterogeneous Signals (Model 3)

One feature of the forecaster data that the model so far does not speak to is the fact that there is heterogeneity in forecasts. A second aspect of forecasting is that obviously, forecasters have access to additional data, beyond just the past realizations of GDP that we allow the forecaster in our models to observe. Additional data about GDP might take the form of leading indicators or firm announcements of future hiring and investment plans.

To model this additional information and the heterogeneity of forecasts, we consider a setting where multiple forecasters update beliefs as in the previous model with state and parameter uncertainty. But each period, each forecaster i observes an additional signal z_{it} that is the next period's GDP growth, with common signal noise and idiosyncratic signal noise:

$$z_{it} = y_{t+1} + \eta_t + \epsilon_{it} \quad (8)$$

where $\eta_t \sim N(0, \sigma_\eta^2)$ is common to all forecasters and $\epsilon_{it} \sim N(0, \sigma_\epsilon^2)$ is i.i.d. across forecasters.

We calibrate the two signal noise variances σ_η^2 and σ_ϵ^2 to match two moments. The idiosyncratic signal noise is chosen to match the average dispersion of forecasts. Forecast dispersion in time t is

$$D_t = \sqrt{\frac{1}{N_t} \sum_{i \in I_t} (E_{it}[y_{t+1}] - \bar{E}_t)^2} \quad (9)$$

where $\bar{E}_t = 1/N_t \sum_i E_{it}[y_{t+1}]$ is the average time- t growth forecast and N_t is the number of forecasters making forecasts at time t . So, σ_ϵ is chosen so that $1/T \sum_t D_t$ are equal in the model and in the forecaster data. The value of that average dispersion is 1.6%.

The common signal noise is chosen to match the average forecast error. Specifically, given a σ_ϵ , we choose σ_η so that the average forecast error, $1/T \sum_t FE_t$ is identical in the model and in the data. The value of that average forecast error is 1.91%. Note that this average error is larger than the dispersion. It tells us that signal noise is not exclusively private signal noise.

Our agent uses the following updating equation to form forecasts. (See appendix B.2 for details)

$$E(y_{t+1}|y^t, z_i^t) = \sum_{S_{t+1}} \int y_{t+1} p(y_{t+1}|\theta, S_{t+1}) \Pr(S_{t+1}|\theta, y^t, z_i^t) p(\theta|y^t, z_i^t) d\theta. \quad (10)$$

Comparing model 3 to forecast data While the model with heterogeneous signal allows us to talk about the part of forecast errors that come from forecast dispersion, the amount of dispersion does not drastically alter our predictions. In a model with must less dispersion (0.4%), we find almost identical levels of uncertainty and correlations of uncertainty with GDP, but larger shocks to uncertainty.

But the fact that forecasters have signals about GDP growth makes a big difference. This model does a better job than any of the others in matching the average forecast error. But recall that we calibrated signal noise in order to match this moment of the data. Reducing the average error also brings down the standard deviation of forecast errors, to a level just slightly above that in the data.

The biggest success of the signal model is that it allows forecasts to be more highly

correlated with future GDP growth. The correlation of 0.74 in this model is almost equal to the correlation of 0.72 in the data. While the other forecasting models do not have enormous signal errors because they are close to the true GDP growth number on average, they miss lots of the little ups and downs in GDP and therefore achieve low correlation. With a signal about future GDP, forecasters in this model are more likely to revise their forecasts slightly up when growth will be high and down when it will be low, achieving a higher correlation between forecast and GDP growth. But the bottom line is that building forecast heterogeneity in the model, similar to what we see in the data, is not a mechanism for generating large uncertainty shocks.

Should signal precision vary over time? One potential source of uncertainty shocks could be changes in the precision of signal. Here, we argue that this is an unlikely source of uncertainty shocks because it suggests other features of the data that are counter-factual. Suppose that in periods where the variance of the noise σ_η or σ_ϵ is high, z_t is a relatively poor predictor of y_{t+1} . Since agents' signal about y_{t+1} are low-precision, their uncertainty about y_{t+1} , conditional on this signal, will be high.

Of course, if all else is equal, a more volatile e_t would mean a more volatile y_t , and this story of forecasting variable precision changes would be the same as a story where uncertainty shocks come from volatility shocks. But it is possible that a fall in signal precision (increase in $\text{var}(e_t)$) could generate an uncertainty shock without a volatility shock to the GDP process y_t . For example, if z and e are independent, then $\text{var}(y_t) = \text{var}(z_t) + \text{var}(e_t)$, and a negative relationship between $\text{var}(z_t)$ and $\text{var}(e_t)$ could leave $\text{var}(y_t)$ unchanged. But this structure would imply that more uncertainty from less informative signals (high $\text{var}(e_t)$) is associated with lower macro volatility (low $\text{var}(z_t)$). There is no such inverse relationship between uncertainty and the volatility of forecasting variables in the data.

There is another way that $\text{var}(y)$ could be constant, despite a shock to signal precision. If e_t and z_t themselves are negatively correlated, then $\text{var}(z)$ and $\text{var}(e)$ can both rise. Since $\text{var}(y) = \text{var}(z) + \text{var}(e) - 2\text{cov}(y, e)$, then a sufficiently high covariance will allow signal precision $1/\text{var}(e)$ to change, resulting in an uncertainty shock, without a volatility shock. But the problem with this explanation is that then z is no longer an unbiased forecast of y . If we transform z to make it an unbiased signal, then this negative correlation with the estimation error would disappear. Exploring these possibilities makes the point that

it is hard to see how changes in the precision of forecast variables can explain uncertainty shocks, in a way that is consistent with the data.

1.6 Considering Fat-Tailed Shocks (Model 4)

One feature of our model so far is that all the innovations are normal. When thinking about the variation in uncertainty, this could be an important feature because normal distributions have the property that the conditional variance is independent of the conditional mean. So, no matter where in the distribution, uncertainty depends only on the precision of the information you've seen, not on the content of that information. Considering non-normal distributions opens up the possibility of an additional source of uncertainty coming from the different conditional variance at different places in the distribution. There are many possible non-normal distributions. But one salient feature of the GDP data is its excess kurtosis. The histogram of GDP data reveals that GDP growth is similar to a normal distribution, but fatter-tailed. A distribution that fits the data as well as a normal, with fatter tails and with a density that can be easily computed is the t-distribution. We do not want to force our forecaster to believe in a model with fat tails. Rather, we introduce model uncertainty. We allow the forecaster to assess the probability that the true data-generating process has either normal shocks or t-distributed shocks and to update this probability each period, as new data is observed.

Of course, there are an infinite number of t-distributed shock distributions, each indexed by a different degree of freedom. To keep the computation tractable, we allow the forecaster to consider a normal shock distribution and one t distribution, fixing the degrees of freedom. The results here use 4 degrees of freedom, the maximum likelihood estimate. This gives a substantial amount of excess kurtosis.

How does the forecaster update his beliefs when he considers multiple generating processes for the observed data series $\{y_t\}$? Let $p(y_{t+1}|y^t, \mathcal{M}_k)$ denote the conditional density of next-period values of y_t conditional on a particular model being contemplated, $\mathcal{M}_k$, and the observed history, y^t . The agent assigns a prior probability $p(\mathcal{M}_k)$ to each of the models, and computes the posterior model weight using Bayes' rule:

$$p(\mathcal{M}_k|y^t, \theta) = \frac{p(y^t|\mathcal{M}_k, \theta) p(\mathcal{M}_k)}{\sum_{m=1}^K p(y^t|\mathcal{M}_m, \theta) p(\mathcal{M}_m)} \quad (11)$$

where the likelihood of the sequence, $p(y^t|\mathcal{M}_k)$, is the product of the conditional likelihoods of each observation

$$p(y^t|\mathcal{M}_k, \theta) = \prod_{i=0}^t p(y_i|y^{i-1}, \mathcal{M}_k, \theta)$$

Bayes' rule ensures that the model probabilities sum to one: $\sum_{k=1}^K p(\mathcal{M}_k|y^t) = 1$. Given the model probability, and the conditional density of y , the probability density of next-period y outcomes $p(y_{t+1}|y^t, \theta, \mathcal{M}_k)$ is $\sum_{k=1}^K p(\mathcal{M}_k|y^t, \theta) p(y_{t+1}|y^t, \mathcal{M}_k, \theta)$. Then, the forecast is

$$E(y_{t+1}|y^t) = \sum_{k=1}^K p(\mathcal{M}_k|y^t) \int y_{t+1} p(y_{t+1}|y^t, \mathcal{M}_k, \theta) p(\theta|y^t) d\theta. \quad (12)$$

Comparing model 4 to forecast data Column 5 of table 2 reveals that this model with model uncertainty and the possibility of fat-tailed shocks fits the observable features of forecast data about equally well as the model with normal shocks. It has forecast errors that are slightly more autocorrelated than the previous models and more correlated than the data. Updating model probabilities introduces even more persistence in the learning process.

1.7 Model-Implied Uncertainty

One striking feature of our analysis is that, no matter how we alter our forecasting model, we get very consistent predictions for the time series of forecast errors. Figure 1 shows the time series plots of the forecast errors, $|E[y_{t+1}|y^t] - y_{t+1}|$ for each of the five models we consider. For the model with heterogeneous forecasts (labeled 'signals'), the line depicts the average forecaster's error in period t . Much of the time, the series are indistinguishable.

Another feature is that uncertainty varies quite a bit. Figure 2 plots the time series of the conditional variances of forecasts, $Var[y_{t+1}|y^t]$ for each of the five models we consider. There is consistent evidence of fluctuations in the model-implied uncertainty, particularly around recessions.

Adding signals to the model clearly lowers uncertainty and smooths it out. The signals are an additional source of information that reduces the conditional variance of agents' estimates.

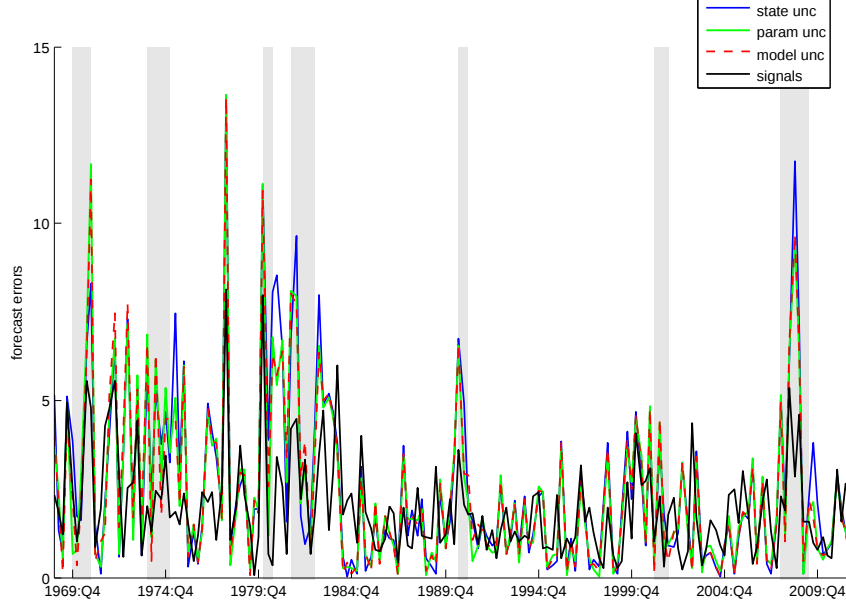


Figure 1: Forecast errors in each of the models.

The remaining three models are more similar. At the start of the sample, the presence of unknown parameters raises uncertainty above what the model with state uncertainty alone predicts. It does so directly, but also indirectly, as the unknown parameters make the state harder to infer. But by the end of the sample, beliefs about parameters have largely converged and the uncertainty levels are similar. Introducing the possibility of fat-tailed innovations ('model unc') achieves the largest fluctuation in uncertainty of the three models.

model	state unc (1)	param unc (2)	signals (3)	model unc (4)
Avg conditional std	3.34%	3.61%	2.63%	3.18%
Std conditional std	0.32%	0.53%	0.23%	0.84%
CoefVar conditional std	0.10	0.15	0.09	0.27
Autocorr conditional std	0.62	0.90	0.89	0.92
Corr(cond std_t , GDP_t)	-0.29	-0.20	-0.15	-0.19
Corr(cond std_t , GDP_{t+1})	-0.19	-0.04	-0.04	-0.04

Table 3: **Properties model uncertainty series.** Column 1 labeled 'state unc' uses equation (4) to forecast y_{t+1} . Columns 2, 3 and 4 use equations (7), (10) and (12) to forecast y_{t+1} .

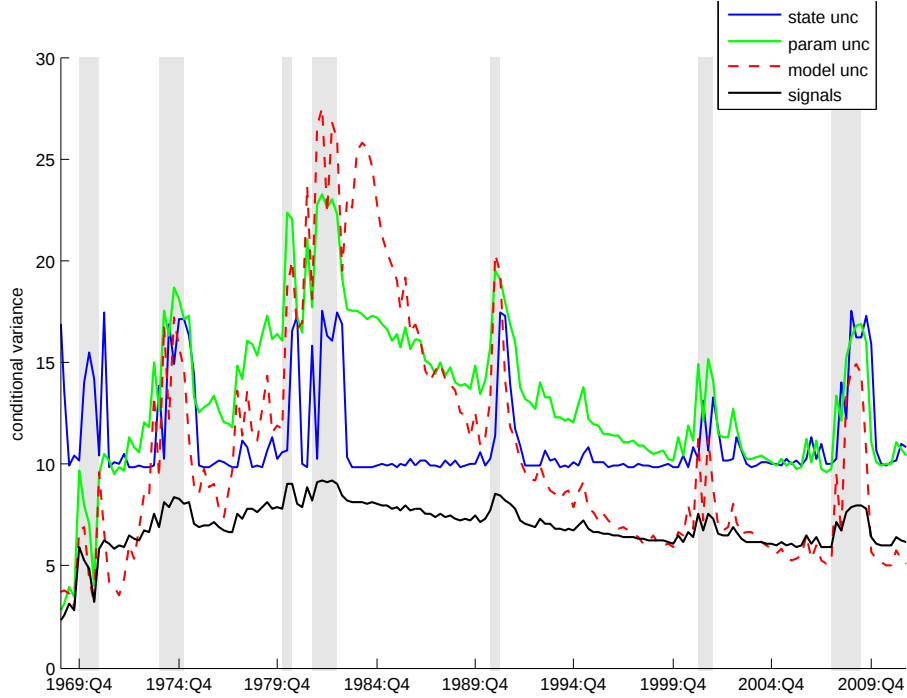


Figure 2: Uncertainty implied by each of the models.

Despite these differences, one important feature that all the models share is that uncertainty is a very persistent process, with low-frequency changes, but not much fluctuation at the business cycle frequency. The persistent uncertainty process comes from the nature of learning: A single large shock to GDP growth results in a quick reevaluation of the parameter and model probabilities. These revisions in beliefs act as permanent, non-stationary shocks even when the underlying shock is transitory.

2 Data Used to Proxy for Uncertainty

Our model generates an endogenous uncertainty series. Next, we'd like to compare our measure to some of the empirical proxies for uncertainty that are commonly used, including the VIX, forecast dispersion, and GARCH volatility estimates. Although these measures are all supposed to serve as proxies for uncertainty, they have different properties from each other and from the model-generated uncertainty series.

2.1 Estimating time-varying volatility

One common procedure for estimating the size of volatility shocks is to estimate an ARMA process that allows for stochastic volatility. In order to compare such a volatility measure to our uncertainty series, we estimate a GARCH model of GDP growth, that allows for time-variation in the variance of the innovations. All data is quarterly and all of these series are non-stationary. To obtain stationary series, we use growth rate.³ We consider annualized rates.

We estimate volatility using an ARMA model with ARCH or GARCH errors. The estimation procedure is maximum likelihood. The choice of each ARMA model is informed by the autocorrelations and partial autocorrelations in the data. We considered several models based on this and chose the AR and MA orders based on the significance of additional variables and their effect on the log-likelihood. Similarly, we considered different lags of linear terms for ϵ_t and variances σ_t^2 in the GARCH specification and used the significance and effect of additional variables on the log-likelihood to inform the specification choice. We assume that the errors in our models, ϵ_t , are Gaussian. We also estimate homoskedastic models for each of the time series to test the hypothesis of homoskedasticity. We use the ARCH-LM test to do this.

The GARCH process that generates the best fit is one with an AR(1) process for GDP growth and a GARCH(1) process for volatility, which includes 1 lagged variance:

$$\Delta gdp_{t+1} = 3.38 + 0.41\Delta gdp_t + \epsilon_{t+1}, \quad \epsilon_{t+1} \sim \mathcal{N}(0, \sigma_{t+1}^2) \quad (13)$$

$$\sigma_{t+1}^2 = 0.52 + 0.76\sigma_t^2 + 0.24\epsilon_t^2 \quad (14)$$

This process is estimated on the full post-war data sample (1947-2011). We estimate heteroskedastic models for similar samples for our other variables. Since we will compare these estimates to data from the survey of professional forecasters, which is available only post-1968, we re-estimate each process, using data from the same quarters as the survey data. The complete set of estimates, as well as more detail about the model selection process are reported in appendix A.

³We check stationarity using Augmented Dickey-Fuller Tests. The results are presented in Appendix A. We also estimate a stochastic volatility model, with very similar results.

This estimated GARCH process produces an estimate of the shock variance σ_{t+1}^2 in each period. This is what we call the volatility process.

The evidence for time-varying volatility is weak. The log likelihood of the highest-likelihood heteroskedastic model is only 3% higher than the best-fitting homoskedastic model. With an ARCH-LM test, one cannot reject the null hypothesis of homoskedasticity (pvalue is 0.22). If one then selects the best-fitting heteroskedastic model, takes the realized time-series of GDP data and feeds it into the model to forecast the volatility of next-period innovations, that volatility varies much less than forecast errors. Specifically, the coefficient of variation of MSEs is 50% higher than the coefficient of variation of GDP volatility. Even once we give volatility its best shot, volatility shocks do not look like forecast error shocks.

Robustness checks Throughout the volatility analysis, we perform the following three robustness checks of our results. First, we relax the distributional assumptions. Our baseline analysis assumes that errors ϵ_t have a Gaussian distribution. Using distributions with fatter tails (student-t) yields no difference in estimations or significance. Second, we explore further lags of all variables; either coefficients were not significantly different from zero or the log-likelihood was reduced. Third, we included different lags of linear terms for ϵ_t and variances σ_t^2 in the GARCH specification. Again, the estimated parameters were not significant or the log-likelihood was reduced.

2.2 Measuring forecast errors and dispersion

We consider two uncertainty proxies constructed from the survey of professional forecasters. One is the forecast mean-squared error (MSE_{t+1}) and the other is forecast dispersion. As described in section 1.1, a forecast of period- t GDP growth made at the start of period t , by forecaster i is denoted $E_{it}[y_{t+1}]$. We define a forecast mean-squared error in quarter t as the square root of the average squared distance between the forecast and the realized value

$$MSE_{t+1} = \sqrt{\frac{\sum_{i \in I_t} (E_{it}[y_{t+1}] - y_{t+1})^2}{N_t}}. \quad (15)$$

Some authors (e.g. Baker, Bloom, and Davis (2012) or Diether, Malloy, and Scherbina (2002)) use forecast dispersion (D_t in equation 9). Dispersion and MSE are closely related,

both statistically and theoretically. The difference depends on private or public nature of information. Imprecise private information generates forecast dispersion, while imprecise public information typically does not. Therefore, we decompose these forecast errors into their public and private components.

The idea that forecast errors are related to an increase in dispersion is not incompatible with our approach. In fact, dispersion in forecasts is one source of mean-squared forecasting error. To see this, note that our uncertainty measure is a square root of the average over all agents j of the expected value of $FE_{jt}^2 = (E_{jt}[y_{t+1}] - y_{t+1})^2$. We can split up FE_{jt}^2 into the sum $((E_{jt}[y_{t+1}] - \bar{E}_t[y_{t+1}]) + (\bar{E}_t[y_{t+1}] - y_{t+1}))^2$, where $\bar{E}_t[y_{t+1}] = \int_j E_{jt}[y_{t+1}]$ is the average forecast. If the first term in parentheses is orthogonal to the second, $1/N \sum_j FE_{jt}^2 = MSE_t^2$ is simply the sum of forecast dispersion and the squared error in the average forecast: $(E_{jt}[y_{t+1}] - \bar{E}_t[y_{t+1}])^2 + (\bar{E}_t[y_{t+1}] - y_{t+1})^2$.

This then raises the question, how much of the variation in mean-squared errors (MSE) comes from changes in the accuracy of average forecasts and how much comes from changes in dispersion? For each macro variable we examine, at least half of the variation comes from average forecasts, but the exact amount depends on the macro variable.

	GDP	FedSpend	SLSpend	IntRt
R^2	80%	54%	97%	65%

Table 4: **The fraction of MSE variation explained by average forecast errors.** This is the R^2 of a regression of the forecast mean-squared error, MSE_t^2 , defined in (15) on $(\bar{E}_t[y_{t+1}] - y_{t+1})^2$. The remaining variation is due to changes in forecast dispersion.

Forecast bias Recent research argues that professional forecasts are biased. (See Elliott and Timmermann (2008) for a review.) While this literature has argued that stock analysts over-estimate earnings growth and the Federal Reserve under-estimates GDP growth, none of it argues that the bias is volatile. In other words, forecast bias may well exist and may explain why forecasts are further away from the true outcome than they perhaps should be. But uncertainty shocks come from fluctuations in the expected size of forecast errors. A fixed bias does not create fluctuations in forecast errors.

How much fluctuation in forecast errors comes from small samples? Another question arises about whether the fluctuations in forecast errors or dispersion could arise simply because of the small sample of forecasts and measurement error or random forecast error. To address this, we create an artificial panel of forecasts with the same number of forecasters on average and homoskedastic error. We then compute MSEs on this artificial panel. Of course, there is some variation period-to-period in MSEs. But that variation is only 0.40, about two-thirds of what is observed in the data. In fact, among 10,000 runs of this artificial process, not one produced variation in uncertainty that is as large as what we see in the data. Even when we create a realistic panel of forecasts about the actual GDP data, we only get a coefficient of variation of MSE's of 0.41, instead of 0.40. This fluctuation in forecast errors does seem to capture some time-variation in the ability to forecast accurately, not just random forecast noise.⁴

2.3 The market volatility index

We discuss the volatility and forecast-based measures in most detail because there are measures that we can relate explicitly to our theory. Other proxy variables for uncertainty are interesting and worth of analysis but have a less clear connection to our model. The market volatility index (VIX) is a traded blend of options that measures expected percentage changes of the S&P500 in the next 30 days. It captures expected volatility of equity prices. When the macroeconomy is very uncertain, more volatile equity prices are probably more likely. But it would take a rich and complicated model to link macroeconomic uncertainty to precise movements in the VIX. Nevertheless, we can compare its statistical properties to those of the uncertainty measure in our model. Figure 3 does just this.

Another commonly cited measure of uncertainty is business or consumer confidence. The consumer confidence survey asks respondents whether their outlook on future business or employment conditions is “positive, negative or neutral.” Likewise, the index of consumer

⁴Our simulated GDP growth process is a series of i.i.d. normal variables with mean 1 and the same variance as true GDP growth. Signals are GDP growth plus two error terms $e + \eta$. Both errors are mean-zero, normal random variables. e is common to all forecasters and η is independent across forecasts. Forecasts are the posterior expectation of GDP growth. The variance of e and η are chosen to match the average mean-squared error and the average dispersion in forecasts, 0.66% and 0.39% respectively. Simulating 10,000 runs of 172 quarters each, with 35 forecasters, we found that the coefficient of variation of uncertainty averaged 0.40, with a minimum of 0.31, a maximum of 0.50, and a standard deviation of 0.026.

sentiment asks respondents whether future business conditions and personal finances will be “better, worse or about the same.” While these indices are indeed negatively correlated with the GARCH-implied volatility of GDP, they are not explicitly questions about uncertainty. Furthermore, we would like to use a measure that we can compare to the forecasts in our model. Since it is not clear what macro variable “business conditions” or “personal finances” corresponds to, it is not obvious what these macro variable respondents are predicting.

2.4 Comparing uncertainty proxies to model-generated uncertainty

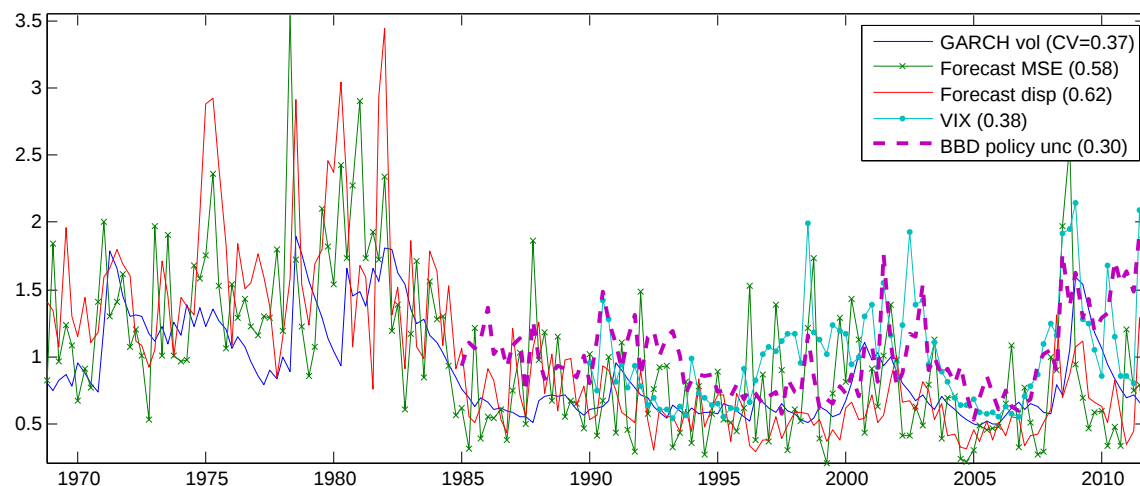


Figure 3: Comparing variables used to measure uncertainty in the literature.

Figure 3 plots each of the uncertainty proxies. There is considerable comovement, but also substantial variation in the dynamics of each process. These are clearly not measures of the same stochastic process, each with independent observation noise. Furthermore, they have properties that are quite different from our model-implied uncertainty metric. Table 5 shows that our uncertainty metric is less volatile, less counter-cyclical, but more persistent than the proxy variables. When we use a simple linear regression to determine which proxies best explain our model-based uncertainty, we find that the only specifications that include variables whose coefficients are significantly different from zero are a specification with a constant and forecast MSE’s and another specification with a constant

	Mean	Coeff Var (std/mean)	autocorr	correlation with g_{t+1}	Sample Period
forecast MSE	2.64%	0.58	0.48	0.04	1968:Q4-2011:Q4
forecast dispersion	1.54%	0.62	0.74	-0.19	1968:Q4-2011:Q4
GARCH volatility	3.65%	0.37	0.90	0.06	1947:Q2-2011:Q4
GARCH real-time	3.45%	0.30	0.90	-0.07	1947:Q2-2012:Q2
VIX	20.55	0.38	0.58	-0.41	1990:Q1:2011:Q4
BBD policy uncertainty	105.95	0.30	0.65	-0.41	1985:Q1:2011:Q4
model 4 uncertainty	3.18%	0.27	0.92	-0.02	1968:Q4-2011:Q4

Table 5: **Properties of forecast errors and volatility series for macro variables.** Forecast MSE and dispersion are defined in (15) and (9). Growth forecast is constructed as $\ln(E_t(GDP_t)) - \ln(E_t(GDP_{t-1}))$. GARCH volatility is the σ_{t+1} that comes from estimating (14) on the full sample of data. At each date t , GARCH real-time is the σ_{t+1} that comes from estimating (14) using only data from 1947:Q2 through date t . VIX_t is the Chicago Board Options Exchange Volatility Index closing price on the last day of quarter t . BBD policy uncertainty is the Baker, Bloom, and Davis (2012) economic policy uncertainty index for the last month of quarter t .

and forecast dispersion. But since forecast MSE's and forecast dispersion are highly correlated, the specification with both variables finds neither significant. The R-squared of these regressions is quite low (6-7%). Thus, forecast errors and dispersion explain some of the variation in uncertainty, but not a lot of it.

2.5 Is uncertainty countercyclical?

Many authors argue that times of high uncertainty trigger recessions. Thus, it is useful to see how our measure of uncertainty and other uncertainty proxies are related to GDP. To be sure, many measures of uncertainty and volatility are higher in recessions than in expansions. But bear in mind that recessions, periods of negative growth, are outliers, where growth is more than 2% below average. This raises the question, is high uncertainty associated with all unusual events, or primarily with low-growth events?

To answer this question, we do the typical breakdown of uncertainty measures in positive and negative growth periods. But then we do a similar split in above-median and below-median growth periods. Table 6 shows that most of the heightened uncertainty arises from outlier events. Lower than median growth raises uncertainty by much less. When we extend the GARCH analysis to the entire post-war sample, the relationship between

higher-than-median growth and uncertainty flips. Higher growth periods turn out to have higher average GARCH-measured volatility.

	Model 4 uncertainty	Professional Forecaster MSE	GARCH Volatility	GARCH Volatility ('47-'12)
Average conditional on positive GDP growth	1.34%	2.44%	3.13%	3.65%
Average conditional on negative GDP growth	1.63%	3.72%	4.00%	4.26%
Average conditional on above median GDP growth	1.33%	2.68%	3.13%	3.78%
Average conditional on below median GDP growth	1.44%	2.56%	3.38%	3.71%

Table 6: Uncertainty in high- and low-growth quarters.

Model, forecast MSEs and first volatility estimates come from data from 1968:Q4–2011:Q4 . Forecast mean-squared errors (MSE) are computed on survey of professional forecaster data, using (15). Second volatility column uses data from 1947q2 – 2012q2. Average model-implied uncertainty comes from model 4, with parameter and model uncertainty.

Of course, if uncertainty causes GDP growth to fall, then perhaps there should be a negative relationship between uncertainty this quarter and output over the following quarter. Table 7 describes the correlation of GDP uncertainty and volatility with GDP leading by -2, -1, 0, 1 and 2 quarters.

	GDP	GDP lead 1	GDP lead 2	GDP lag 1	GDP lag 2
Model 4	-0.02	0.011	0.05	-0.17	-0.08
Forecast MSE	0.0441	-0.1080	-0.0446	-0.0863	-0.0901
Volatility	-0.1022	-0.0681	-0.0356	-0.0987	-0.1100

Table 7: Correlation of uncertainty and volatility with GDP growth and GDP growth leading and lagging by 1 and 2 quarters, 1968:Q4–2011:Q4.

3 Considering Policy Uncertainty

One aspect of uncertainty that our analysis so far has neglected is policy uncertainty. Perhaps the discrepancy between volatility and uncertainty can be explained by changes in the uncertainty about policy shocks. We consider two approaches to this comment. The first is to compare the properties of the Baker, Bloom, and Davis (2012) policy uncertainty

index to our GDP uncertainty measures, just like we do with all the other proxies for uncertainty. The second approach is to use the model to help think about the question: Where do policy uncertainty shocks come from?

Table 5 describes the summary statistics of the BBD policy uncertainty index. This index is a mix of legal changes, textual analysis, forecast dispersion. The nature of some of this data makes it impossible to derive the index from the model. But what we can say is that this index has shocks that are about as large in relative magnitude as our uncertainty measure. The coefficient of variation is 0.30 for the index and 0.27 for the model. But the model uncertainty is much more persistent (0.92 vs. 0.65) and less counter-cyclical (-0.02 vs -0.41 correlation with GDP growth).

		Mean	Coeff Var (std/mean)	Sample Period
Fed Govt	forecast MSE	6.36%	0.53	1981:Q3-2012:Q1
Spending	volatility	5.90%	0.01	1981:Q3-2012:Q1
S&L Govt	forecast MSE	6.60%	0.62	1981:Q3-2012:Q1
Spending	volatility	2.36%	0.268	1981:Q3-2012:Q1
Interest	forecast MSE	1.00%	0.82	1981:Q3-2012:Q1
Rate	volatility	0.47%	0.65	1981:Q3-2012:Q1

Table 8: **Properties of forecast errors and volatility series for macro variables.** All variables are in growth rates, except for the interest rate, which is the first difference of the level. So a 1% mean for uncertainty means that the average forecast of the quarterly rate of growth has a 1% standard deviation, relative to the final reported value. Uncertainty denotes the square root of the mean squared error of current period growth forecasts. This growth forecast is constructed as $\ln(E_t(x_t)) - \ln(E_t(x_{t-1}))$.

Our second approach to thinking about policy uncertainty is to point out that there is a potential role for our model in explaining policy uncertainty shocks. Table 8 reports the moments of volatility and forecast errors for three policy variables: federal spending, state and local spending and interest rates. For the matching subsample of data, state and local spending has forecast MSE shocks that are 3.4 times larger (coefficient of variation is 0.62) than the shocks to volatility (coefficient of variation is 0.18). For federal spending, the differences in coefficients of variation is even larger (0.53 vs 0.02). The forecast MSE shocks are 26 times larger than the volatility shocks. This large gap between volatility shocks and forecast error shocks could have a couple of possible explanations. One is

that forecasts contain lots of public information whose quality is changing dramatically over time. Another possibility is that model and parameter uncertainty explains this gap. Estimation errors in model and parameters create time-variation in forecast errors, but not volatility.

For the interest rate, volatility and forecast errors are not nearly so dissimilar. Forecast MSE shocks are 15% larger than volatility shocks. But what is obvious for interest rates is that the average size of forecast errors is very small. In other words, agents have very precise information about future interest rates. Of course, interest rates are lower in level than the size of government spending. But, keep in mind that all series are log-differenced. So, we are comparing percentage (log) changes in spending and GDP to changes in interest rates, which makes them comparable. Thus, the low level of uncertainty and volatility really reflects the fact that the degree of uncertainty about the other variables is much greater.

The conclusion we draw from this discrepancy is that, just like GDP uncertainty shocks do not seem to come primarily from volatility shocks, policy uncertainty shocks also do not seem to be driven by changes in policy volatility. Rather, they are a product of the learning process, which this paper explores.

4 Inferring Uncertainty about Productivity

Most uncertainty shock business cycle models have fluctuations in uncertainty about aggregate productivity as their driving process. So what does TFP uncertainty look like? We now estimate our model to TFP data to infer a time-series of TFP uncertainty shocks.

We use model 2 (only state and parameter uncertainty) to infer a stochastic process for TFP uncertainty. Specifically, we estimate equation (7) using TFP data from 1947-2011. Then, to see how considering fat tails changes our predictions, we estimate a second model, again with no model uncertainty, but where we assume the innovations are t-distributed with 5 degrees of freedom. The resulting series of uncertainties are plotted in figure 4. These stochastic processes could then be used by others to fit their models.

The salient result here is the large fluctuations in uncertainty, particularly in the model with fat tails. The short run changes are large, there is a clear cyclical pattern, and there is a salient low-frequency component as well.

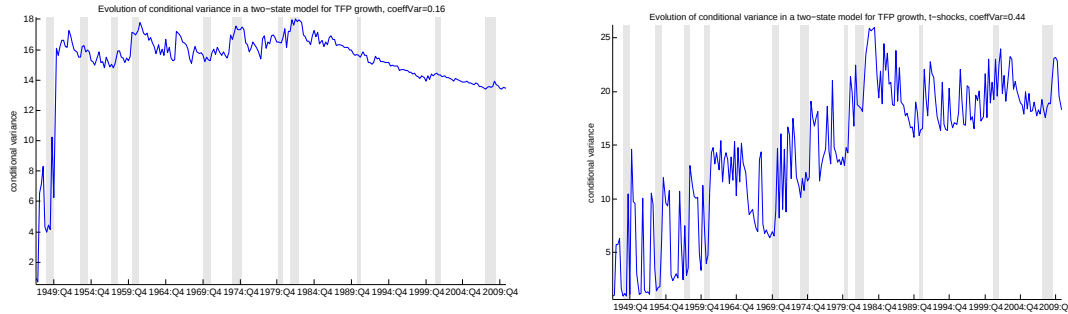


Figure 4: TFP uncertainty from a model with normal innovations (left) and fat-tailed innovations (right). The complete set of parameters for both models is listed in table 5.

5 Conclusions

The data typically used to measure uncertainty offers a muddy picture of what uncertainty shocks look like: Survey data reveals large swings in the ability of agents to make accurate forecasts. Yet, there is weak evidence of time-varying variance of macro aggregates and the estimated magnitude of these shocks is much smaller than what would explain the survey data. Uncertainty shocks appear to come not from the properties of the data being forecast, but from some state of mind of the forecaster himself.

Our model reconciles these findings. It takes the actual series of GDP, with whatever heteroskedasticity it has, as an input and generates forecasts with properties similar to those in the forecast data. But with this model structure, we can compute actual uncertainty – the conditional variance of the forecast of next quarter growth. We find that this uncertainty process looks different from any of the proxy variables, with smaller shocks and more persistence.

The model also helps us to understand where uncertainty shocks come from. People use simple models to forecast complex economic processes. They have to. Any model that is as complicated as the economy itself would be intractable and not useful. Recognizing this, it is natural for agents to consider more than one possible model of the world. Of course forecasters also are not endowed with knowledge of the parameters of their model. They estimate them over time. When we estimate a forecasting model with these three ingredients: state uncertainty, model uncertainty and parameter uncertainty, we find that even a homoskedastic process can generate time-varying uncertainty that resembles what

we find in forecast data.

A next step for this agenda is to explore firm-level earnings forecasts and firm-level shocks and determine whether our model can also explain uncertainty facts at the micro level.

References

- BACHMANN, R., AND C. BAYER (2012a): “Investment Dispersion and the Business Cycle,” University of Aachen working paper.
- (2012b): “Wait-and-See Business Cycles?,” University of Aachen working paper.
- BACHMANN, R., S. ELSTNER, AND E. SIMS (2012): “Uncertainty and Economic Activity: Evidence from Business Survey Data,” University of Aachen working paper.
- BAKER, S., N. BLOOM, AND S. DAVIS (2012): “Measuring Economic Policy Uncertainty,” Stanford University working paper.
- BANSAL, R., AND I. SHALIASTOVICH (2010): “Confidence Risk and Asset Prices,” *American Economic Review*, 100(2), 537–41.
- BASU, S., AND B. BUNDICK (2012): “Uncertainty Shocks in a Model of Effective Demand,” Boston College working paper.
- BIANCHI, F., C. ILUT, AND M. SCHNEIDER (2012): “Ambiguous business cycles and uncertainty shocks,” Stanford University working paper.
- BLOOM, N. (2009): “The Impact of Uncertainty Shocks,” *Econometrica*, 77(3), 623–685.
- BLOOM, N., M. FLOETOTTO, N. JAIMOVICH, I. SAPORA-EKSTEN, AND S. TERRY (2012): “Really Uncertain Business Cycles,” NBER working paper 13385.
- BORN, B., A. PETER, AND J. PFEIFER (2011): “Fiscal News and Macroeconomic Volatility,” Bonn econ discussion papers, University of Bonn, Germany.
- BRUNO, V., AND H. SHIN (2012): “Capital Flows, Cross-Border Banking and Global Liquidity,” Princeton University working paper.
- CHRISTIANO, L., R. MOTTO, AND M. ROSTAGNO (2012): “Risk Shocks,” Northwestern University working paper.
- COGLEY, T., AND T. SARGENT (2005): “The Conquest of US Inflation: Learning and Robustness to Model Uncertainty,” *Review of Economic Dynamics*, 8, 528–563.
- DIETHER, K., C. MALLOY, AND A. SCHERBINA (2002): “Differences of Opinion and the Cross-Section of Stock Returns,” *Journal of Finance*, 57, 2113–2141.
- ELLIOTT, G., AND A. TIMMERMANN (2008): “Economic Forecasting,” *Journal of Economic Literature*, 46(1), 3–56.

- FERNANDEZ-VILLAVERDE, J., P. GUERRON-QUINTANA, J. F. RUBIO-RAMIREZ, AND M. URIBE (2011): “Risk Matters: The Real Effects of Volatility Shocks,” *American Economic Review*, 101(6), 2530–61.
- FERNANDEZ-VILLAVERDE, J., K. KUESTER, P. GUERRON, AND J. RUBIO-RAMIREZ (2012): “Fiscal Volatility Shocks and Economic Activity,” University of Pennsylvania working paper.
- HANSEN, L. (2007): “Beliefs, Doubts and Learning: Valuing Macroeconomic Risk,” *American Economic Review*, 97(2), 1–30.
- JOHANNES, M., L. LOCHSTOER, AND Y. MOU (2011): “Learning About Consumption Dynamics,” Columbia University working paper.
- NAKAMURA, E., D. SERGEYEV, AND J. STEINSSON (2012): “Growth-Rate and Uncertainty Shocks in Consumption: Cross-Country Evidence,” Columbia University working paper.
- NIMARK, K. (2012): “Man Bites Dog Business Cycles,” Working paper.
- PASTOR, L., AND P. VERONESI (2012): “Uncertainty about Government Policy and Stock Prices,” *Journal of Finance*, forthcoming.
- ROBERTS, G., A. GELMAN, AND W. GILKS (1997): “Weak Convergence and Optimal Scaling of Random Walk Metropolis Algorithms,” *Annals of Applied Probability*, 7(1), 110–120.

A Estimated ARCH/GARCH process for each macro variable

In this section we present the models that we estimate for TFP, GDP, government spending and the three month treasury rate. For each variable we perform our analysis for two time periods. So that we have volatility estimates which are comparable to the uncertainty estimates presented in Section 2, we perform estimation for the subset of the available data which overlaps with the corresponding uncertainty data. We call these subsets ‘matching subsamples.’ We also perform the analysis on the full period of available data to check that each matching subsample is representative of the full sample from which it is taken. The only exception to using all available data is for the government spending variables. For reasons outlined in section A.0.2, we drop some of the earliest data for these variables.

For each sample period of each variable we use Augmented Dickey-Fuller (ADF) tests to check that the series is stationary. This test is based on estimating an AR model for

the data. For each sample we use the best-fitting AR model as the basis for the test.⁵ The null-hypothesis for the test is that the series has a unit root and the alternative hypothesis is that the series is stationary. Where the data permits, we also perform the test on the levels data (i.e. the data that has not been first-differenced). In this case the tests indicate that the first difference series is stationary if we cannot reject the null hypothesis for the levels data and we can reject the null hypothesis for the first-difference of the data. We have the data to do this for government spending and the interest rate.

As discussed in Section 2, we estimate homoskedastic and heteroskedastic models for each series. For the homoskedastic models we use robust standard errors. We use the errors from the homoskedastic models to test the null hypothesis of no heteroskedasticity using an ARCH-LM test. This test requires fitting an AR model to the squared errors of the homoskedastic model. For each sample we set the AR order equal to the order of the AR component of the ARCH/GARCH model that we use.

A.0.1 GDP

We use quarterly GDP growth rate data for 1947:Q2–2012:Q2.⁶ The subsample of the GDP data which matches the uncertainty data is 1968:Q4 to 2011:Q4. The results of ADF tests (Table 9) indicate that the series for the full sample and matching subsample are stationary.

Best-fitting homoskedastic model (matching subsample) Recall that we defined $y_{t+1} \equiv \ln(gdp_t) - \ln(gdp_{t-1})$. The best-fitting process is an ARMA(1,0):

$$y_{t+1} = 2.81 + 0.31y_t + \epsilon_{t+1}, \quad \epsilon_{t+1} \sim \mathcal{N}(0, 10.95).$$

The log-likelihood is -452.51. All coefficients are significant at 1%. The p-value for the ARCH-LM test is 0.219 so homoskedasticity cannot be rejected.

Best-fitting heteroskedastic model (matching subsample) The best-fitting process is ARMA(1,0) for growth and GARCH(1,1) for variance:

$$y_{t+1} = 3.145 + 0.385y_t + \epsilon_{t+1}, \quad \epsilon_{t+1} \sim \mathcal{N}(0, \sigma_{t+1}^2), \quad \sigma_{t+1}^2 = 0.515 + 0.767\sigma_t^2 + 0.217\epsilon_t^2.$$

⁵Note that the number of lags used for the test is one less than the order of the AR process that we model the data with. If we are using an AR(p) process to model the variable x_t then the regression underlying the ADF test is

$$\Delta x_t = \phi x_{t-1} + \sum_{j=1}^{p-1} \alpha_j \Delta x_{t-j} + u_t, \tag{16}$$

where u_t is the error term for an OLS regression.

⁶Data source: Bureau of Economic Analysis (<http://www.bea.gov/national/index.htm#gdp>). We are using the seasonally adjusted annual rate for the quarterly percentage change in real GDP. Note that this data is for the actual percentage change, not an approximation. The version of the data is 27 July 2012, downloaded 2 August 2012.

The log-likelihood is -438.99. All coefficients except the constant term for the GARCH process are significant at 1%.

A.0.2 Government spending

Our data on federal spending and state and local government spending is for 1968:Q4 to 2012:Q2.⁷ So that we have a stationary series we analyze the annualized quarterly growth rate, approximated using the same method as for TFP. The results of ADF tests (Table 9) are consistent with our growth rate series being stationary.

Best-fitting homoskedastic model (matching subsample) The best-fitting model for federal spending is an ARMA(3,3):

$$\Delta g_{t+1} = 1.759 + 1.450\Delta g_t - 1.283\Delta g_{t-1} + 0.726\Delta g_{t-2} + \varepsilon_{t+1} - 1.578\varepsilon_t + 1.575\varepsilon_{t-1} - 0.761\varepsilon_{t-2},$$

$$\varepsilon_{t+1} \sim \mathcal{N}(0, 41.940).$$

For state and local government spending we use an ARMA(1,1):

$$\Delta g_{t+1} = 1.816 + 0.850\Delta g_t + \varepsilon_{t+1} - 0.485\varepsilon_t, \quad \varepsilon_{t+1} \sim \mathcal{N}(0, 5.806).$$

All parameters are significant at 1% except the constant for federal spending (not significant) and the constant for state and local spending (significant at 5%). The log-likelihoods are -405.42 and -282.95 respectively. The p-values for the ARCH-LM tests are 0.672 and 0.000 respectively. Therefore homoskedasticity can be rejected for state and local government spending, but cannot be rejected for federal spending.

Best-fitting heteroskedastic model (matching subsample) The best-fitting model for federal spending is an ARMA(3,3) for growth and ARCH(1) for variance:

$$\Delta g_{t+1} = 7.031 + 1.513\Delta g_t - 1.235\Delta g_{t-1} + 0.715\Delta g_{t-2} + \varepsilon_{t+1} - 1.746\varepsilon_t + 1.888\varepsilon_{t-1} - 0.963\varepsilon_{t-2},$$

$$\varepsilon_{t+1} \sim \mathcal{N}(0, \sigma_{t+1}^2),$$

$$\sigma_{t+1}^2 = 34.335 + 0.012\varepsilon_t^2.$$

All parameters are significant at 1% except for the ARCH parameter (insignificant). The log-likelihood is -392.77. For state and local government spending, we use an ARMA (1,1) for growth and ARCH(1) for variance (an AR(2) with an ARCH(1) process fits equally well):

$$\Delta g_{t+1} = 1.487 + 0.863\Delta g_t + \varepsilon_{t+1} - 0.450\varepsilon_t, \quad \varepsilon_{t+1} \sim \mathcal{N}(0, \sigma_{t+1}^2),$$

⁷Source: National Income and Product Accounts, Table 3.9.3, Real Government Consumption Expenditures and Gross Investment, Quantity Indexes. Available at <http://www.bea.gov/national/nipaweb/DownSS2.asp>, downloaded 8 August 2012. The data is seasonally adjusted.

$$\sigma_{t+1}^2 = 3.483 + 0.419\varepsilon_t^2.$$

All parameters are significant at 1% except for the constant for the growth process (significant at 10%) and the ARCH parameter (significant at 5%). The log-likelihood is -276.49.

A.0.3 Interest rates

The interest rate that we use is the three month treasury bill rate.⁸ Ideally we would like to use the policy rate (the federal funds rate), but no forecast data is available for this. We have therefore chosen the interest rate for which forecast data is available whose time horizon is closest to the federal funds rate. We have data on this rate for 1954:Q1–2012:Q2.

So that our data is stationary we first difference it. The results of ADF tests (Table 9) sample support this: at the 5% significance level we cannot reject the null hypothesis that the interest rate level data has a unit root and we can reject this null hypothesis for the first-difference of the data. However the outcome of the test for the levels data is sensitive to the order of the AR model that we fit to the data. If we assume that the levels data is modeled by an AR(1) process instead of an AR(2) process then the null hypothesis is rejected at the 5% level.

Best-fitting homoskedastic model (matching subsample) The best-fitting process is an ARMA(1,0) with no constant:

$$\Delta r_t = 0.313\Delta r_{t-1} + \varepsilon_t, \quad \varepsilon_t \sim \mathcal{N}(0, 0.3606).$$

The AR(1) coefficient is significant at the 5% level, the variance estimate is significant at the 5% level. The log-likelihood is -111.86. The p-value for the ARCH-LM test is 0.045 so homoskedasticity can be rejected at 5%.

Best-fitting heteroskedastic model (matching subsample) The best-fitting process is an ARMA(1,0) with no constant and a GARCH(1,1) for variance:

$$\begin{aligned} \Delta r_t &= 0.663\Delta r_{t-1} + \varepsilon_t, \quad \varepsilon_t \sim \mathcal{N}(0, \sigma_t^2), \\ \sigma_t^2 &= 0.016 + 0.737\sigma_{t-1}^2 + 0.162\varepsilon_{t-1}^2 \end{aligned}$$

All coefficients are significant at the 1% level. The log-likelihood is -81.22.

For both the homoskedastic and heteroskedastic processes for the matching subsample we estimated models with higher order AR and MA components and found that the extra coefficients were not significant.

For TFP and GDP, we did not investigate the stationarity of the levels series because the raw data was already in growth rate form.

⁸Data source: Board of Governors (<http://research.stlouisfed.org/fred2/series/DTB3/>). We use the 3-Month Treasury Bill: Secondary Market Rate. The data has a daily frequency and we use the average for each quarter.

	Sample period	No. of lags	Constant (C), Trend (T)	Test stat.	5% crit. value
TFP	1972:Q1–2010:Q4	1	C	-6.516	-2.886
	1947:Q2–2012:Q1	2	C	-9.110	-2.880
GDP	1968:Q4–2011:Q4	0	C	-9.393	-2.885
	1947:Q2–2012:Q2	0	C	-10.930	-2.880
Fed Spending	1981:Q3–2012:Q1	3	C	-3.234	-2.889
	1952:Q1–2012:Q2	3	C	-4.866	-2.881
Fed Spending (level)	1981:Q3–2012:Q1	4	C, T	-1.843	-3.448
	1952:Q1–2012:Q2	4	C, T	-3.645	-3.432
S&L Spending	1981:Q3–2012:Q1	1	C	-4.006	-2.889
	1952:Q1–2012:Q2	2	C	-5.565	-2.881
S&L Spending (level)	1981:Q3–2012:Q1	2	C, T	0.704	-3.447
	1952:Q1–2012:Q2	3	C, T	-1.381	-3.432
Interest Rate	1981:Q3–2012:Q1	0	–	-7.938	-1.950
	1954:Q1–2012:Q2	0	–	-12.005	-2.582
Interest Rate (level)	1981:Q3–2012:Q1	1	C	-1.721	-2.889
	1954:Q1–2012:Q2	1	C	-2.379	-2.881

Table 9: **Results of Augmented Dickey-Fuller tests.** Recall that Fed Spending and S&L Spending have been constructed as $400 \times (\log X_t - \log X_{t-1})$. For the levels series for these variables we therefore use $400 \times \log X_t$. Interest Rate (level) is the series of actual interest rates. In column 4, ‘C’/‘T’ denote that a constant/linear trend term (or both) has been added to regression equation (16).

B Estimating the model

For each of the models with parameter uncertainty (models 2-4), the updating process starts with prior beliefs about each of the model parameters. The complete set of prior means and variances for the model with normally-distributed innovations is in table 10.

B.1 Computational algorithm

In what follows we show how to use Metropolis-Hastings algorithm to generate samples from $p(\theta|y^t)$ for each $t = 0, 1, 2, \dots, T$.⁹

The general idea of MCMC methods is to design a Markov chain whose stationary distribution, π (with $\pi T = \pi$ where T is a transitional kernel), is the distribution p we are seeking to characterize. In particular, the Metropolis-Hastings sampling algorithm constructs an ergodic Markov chain that satisfies a detailed balance property with respect

⁹We drop here the dependence on $\mathcal{M}$ hoping that no confusion arises; the algorithm is applied independently to generate the respective samples under each model.

2-state model			2-state model, fat tails		
Par	Mean	Std	Par	Mean	Std
μ_1	-2.47%	1%	μ_1	-3.02%	1%
μ_2	3.55%	0.7%	μ_2	3.22%	0.8%
σ^2	(2.9%) ²	2.5%	σ^2	(2.17%) ²	2%
q_{11}	0.69	0.16	q_{11}	0.67	0.16
q_{22}	0.96	0.032	q_{22}	0.95	0.05

Table 10: Prior assumptions for normal and t distributed models for GDP growth.

2-state			2-state, t, $\nu = 5$		
Parameter	Mean	Stdev	Parameter	Mean	Stdev
μ_1	-1.55	1	μ_1	0.4	0.23
μ_2	1.92	0.4	μ_2	4.8	0.7
σ^2	(3.56) ²	1.5	σ^2	(2.17) ²	1.6
q_{11}	0.62	0.16	q_{11}	0.86	0.16
q_{22}	0.93	0.06	q_{22}	0.5	0.1

Figure 5: Prior assumptions for normal and t distributed models for TFP growth.

to p and, therefore, produces the respective approximate samples. The transition kernel of that chain, T , is constructed based on sampling from a proposal conditional distribution $q(\theta|\theta^{(d)})$ where d denotes the number of the sampling step. Specifically, given the d -step in the random walk $\theta^{(d)}$ the next-step $\theta^{(d+1)}$ is generated as follows

$$\theta^{(d+1)} = \begin{cases} \theta' & \text{with probability } \alpha(\theta^{(d)}, \theta') = \min\left(1, \frac{p(\theta'|y^t)}{p(\theta^{(d)}|y^t)} \frac{q(\theta^{(d)}|\theta')}{q(\theta'|\theta^{(d)})}\right) \\ \theta^{(d)} & \text{with probability } 1 - \alpha(\theta^{(d)}, \theta') \end{cases}$$

where $\theta' \sim q(\theta|\theta^{(d)})$.

In our application, the simulation of the parameters is done through simple random walk proposals. In particular, for the means the proposed move is

$$\mu'_S = \mu_S + \tau_u \varepsilon_S \quad (17)$$

where $S \in \{1, \dots, N\}$, $\varepsilon_S \sim \mathcal{N}(0, 1)$.

For the variance σ^2 , the proposed move is a multiplicative random walk

$$\log \sigma' = \log \sigma + \tau_\sigma \xi_\sigma \quad (18)$$

where $\xi_\sigma \sim \mathcal{N}(0, 1)$.

In the case of the transition probability matrix, the move is slightly more involved due to the constraint on the sum of rows. We reparameterize each row $(q_{i1}, \dots, q_{iN})$ as

$$q_{ij} = \frac{\omega_{ij}}{\sum_j \omega_{ij}}, \quad \omega_{ij} > 0, \quad j \in \{1, \dots, N\}$$

so that the summation constraint does not hinder the random walk. The proposed move on ω_{ij} is then given by

$$\log \omega'_{ij} = \log \omega_{ij} + \tau_\omega \xi_\omega \quad (19)$$

where $\xi_\omega \sim \mathcal{N}(0, 1)$. Note that this reparametrization requires that we select a prior distribution on ω_{ij} rather than on q_{ij} .

The parameters τ_u , τ_σ , and τ_ω can be adjusted to optimize the performance of the sampler. Choosing a proposal with small variance would result in relatively high acceptance rates but with strongly correlated consecutive samples. See Roberts, Gelman, and Gilks (1997) for the results on optimal scaling of the random walk Metropolis algorithm.

Since the proposal is symmetric in all the cases (17), (18), and (19) and since ε_S , ξ_σ , and ξ_ω are drawn independently from one another, we have $q(\theta|\theta') = q(\theta'|\theta)$, and the acceptance probability simplifies to $\min\left(1, \frac{p(\theta'|y^t)}{p(\theta|y^t)}\right)$. To compute that acceptance ratio, note that the posterior distribution $p(\theta|y^t)$ is given by

$$p(\theta|y^t) = \frac{p(y^t|\theta) p(\theta)}{p(y^t)}$$

where $p(y^t) = \int p(y^t|\theta) p(\theta) d\theta$ is the marginal likelihood (or data density).

The means, medians and confidence intervals of the resulting parameter estimates are illustrated in figure 6 for the model with normally-distributed shocks and in figure 7 for the model with t-distributed shocks.

B.2 Estimating the model with heterogeneous signals

With heterogenous signals, the model is the following

$$\begin{aligned} y_{t+1} &= \mu_{S_{t+1}} + \sigma e_{t+1} \\ z_{it} &= y_{t+1} + \sigma_\epsilon \epsilon_{it} + \sigma_\eta \eta_t \end{aligned}$$

where we assume that η_t , e_t , $\epsilon_{it} \sim N(0, 1)$ and that η_t , and e_t are independent from one another and from ϵ_{it} for all t and all i . The forecaster wishes to compute his forecast

$$E_t(y_{t+1}|y^t, z_i^t) = \int y_{t+1} p(y_{t+1}|y^t, z_i^t) dy_{t+1}$$

where the predictive density is given by

$$\begin{aligned} p(y_{t+1}|y^t, z_i^t) &= \sum_{S_{t+1}} \int p(y_{t+1}|\theta, S_{t+1}) p(\theta, S_{t+1}|y^t, z_i^t) d\theta \\ &= \sum_{S_{t+1}} \int p(y_{t+1}|\theta, S_{t+1}) \Pr(S_{t+1}|\theta, y^t, z_i^t) p(\theta|y^t, z_i^t) d\theta \end{aligned}$$

In other words, our forecaster is still faced with the joint estimation problem of parameters and states, but now has additional information to condition on. We can still use the MCMC methods to approximate the posterior parameter distribution, $p(\theta|y^t, z_i^t)$. To that end, we need to show how to determine the data likelihood for a given parameter vector, $p(y^t, z_i^t|\theta)$, and how the agent updates state beliefs conditional on the observations of GDP growth rates, as well as the signal.

For a given θ , the likelihood $p(y^T, z_i^T)$ can be determined as follows

$$\log p(y^T, z_i^T) = \sum_{t=1}^T \log [p(y_t, z_{it}|y^{t-1}, z_i^{t-1})]$$

In turn,

$$p(y_t, z_{it}|y^{t-1}, z_i^{t-1}) = \sum_{k=1}^2 \sum_{j=1}^2 p(y_t, z_{it}|y_{t-1}, z_{it-1}, S_t = k, S_{t+1} = j) \Pr(S_t = k|y^{t-1}, z_i^{t-1}) \Pr(S_{t+1} = j|S_t = k)$$

where $p(y_t, z_{it}|y_{t-1}, z_{it-1}, S_t = k, S_{t+1} = j)$ is a density of the following multivariate normal distribution

$$MVN \left(\begin{bmatrix} \frac{\mu_{S_t=k} + \frac{z_{it}}{\sigma_\eta^2 + \sigma_\epsilon^2}}{\frac{1}{\sigma^2} + \frac{1}{\sigma_\eta^2 + \sigma_\epsilon^2}} \\ \mu_{S_{t+1}=j} \end{bmatrix}, \begin{bmatrix} \left(\frac{1}{\sigma^2} + \frac{1}{\sigma_\eta^2 + \sigma_\epsilon^2} \right)^{-1} & 0 \\ 0 & \sigma^2 + \sigma_\eta^2 + \sigma_\epsilon^2 \end{bmatrix} \right)$$

To show how the agent updates his beliefs recursively, assume for a moment that we know $\psi_{k,t} \equiv \Pr(S_t = k|g^t, z_i^t)$ with $\sum_{k=1}^2 \psi_{k,t} = 1$. The posterior state belief $\psi_{k,t+1}$ can be updated recursively in the following way

$$\psi_{k,t+1} = \frac{\sum_{m=1}^2 \sum_{j=1}^2 p(y_{t+1}, z_{it+1}|S_{t+1} = k, S_t = j, S_{t+2} = m, y^t, z_i^t) q_{km} q_{jk} \psi_{j,t}}{p(y_{t+1}, z_{it+1}|y^t, z_i^t)}$$

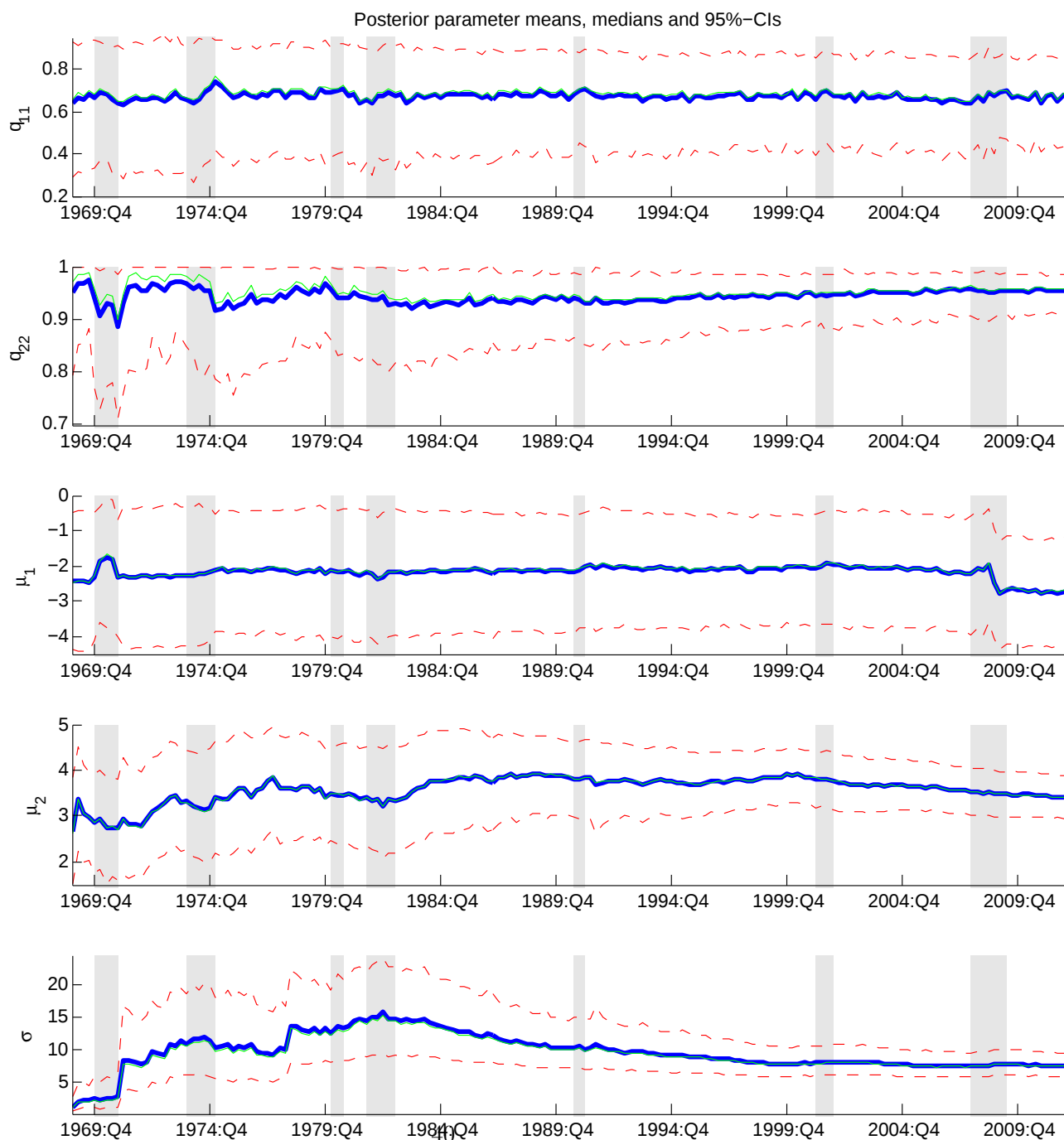


Figure 6: Parameter estimates in model with normal shocks.

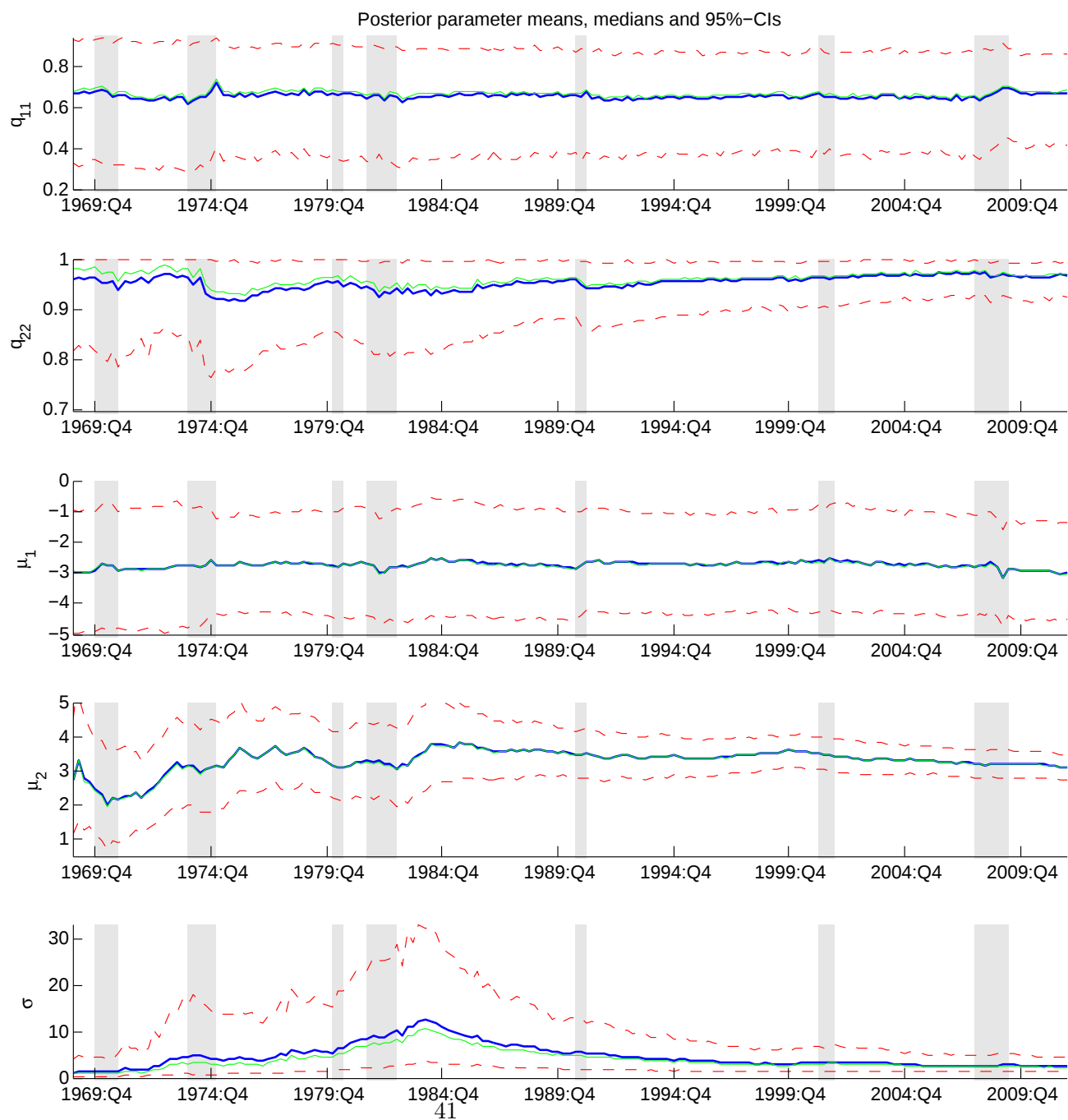


Figure 7: Parameter estimates in model with t -distributed shocks.