

The Zero Lower Bound and Estimation Accuracy*

Tyler Atkinson Alexander W. Richter Nathaniel A. Throckmorton

First Draft: May 7, 2018

This Draft: June 25, 2019

ABSTRACT

During the Great Recession, central banks lowered their policy rate to the zero lower bound (ZLB), calling into question linear estimation methods. There are two alternatives: estimate a nonlinear model that accounts for precautionary savings effects of the ZLB or a piecewise linear model that is faster but ignores the precautionary savings effects. This paper compares their accuracy using artificial datasets. The predictions of the nonlinear model are typically more accurate than the piecewise linear model, but the differences are usually small. There are far larger gains in accuracy from estimating a richer, less misspecified piecewise linear model.

Keywords: Bayesian Estimation; Projection Methods; Particle Filter; OccBin; Inversion Filter

JEL Classifications: C11; C32; C51; E43

*Atkinson and Richter, Research Department, Federal Reserve Bank of Dallas, 2200 N Pearl Street, Dallas, TX 75201 (tyler.atkinson@dal.frb.org; alex.richter@dal.frb.org); Throckmorton, Department of Economics, William & Mary, P.O. Box 8795, Williamsburg, VA 23187 (nat@wm.edu). We thank Boragan Aruoba, Alex Chudik, Marc Gianoni, Matthias Hartmann, Rob Hicks, Ben Johannsen, Giorgio Primiceri, Sanjay Singh, Kei-Mu Yi, and an anonymous referee for suggestions that improved the paper. We also thank Chris Stackpole and Eric Walter for supporting the supercomputers at our institutions. This research was completed with supercomputing resources provided by the Federal Reserve banks of Kansas City and Dallas, William & Mary, Auburn University, the University of Texas at Dallas, the Texas Advanced Computing Center, and Southern Methodist University. The views in this paper are those of the authors and do not necessarily reflect the views of the Federal Reserve Bank of Dallas or the Federal Reserve System.

1 INTRODUCTION

Using Bayesian methods to estimate linear dynamic stochastic general equilibrium (DSGE) models has become common practice in the literature over the last 20 years. Many central banks also use these models for forecasting and counterfactual simulations. The estimation procedure sequentially draws parameters from a proposal distribution, solves the model given that draw, and then evaluates the likelihood function. With linearity and normally distributed shocks, the model solves in a fraction of a second and it is easy to exactly evaluate the likelihood function with a Kalman filter.¹

The financial crisis and subsequent recession compelled many central banks to take unprecedented action to reduce their policy rate to its zero lower bound (ZLB), calling into question linear estimation methods. The ZLB constraint presents a challenge for empirical work because it creates a kink in the central bank's policy rule. The constraint has always existed, but when policy rates were well above zero and the likelihood of hitting the constraint was negligible, it was reasonable to ignore it. The lengthy period of near zero policy rates over the last decade and the increased likelihood of future ZLB events due to estimates of a lower natural rate has forced researchers to think more carefully about the ZLB constraint and its implications (e.g., Laubach and Williams, 2016).

There are two promising estimation methods used in the literature that account for the ZLB constraint in DSGE models. The first method estimates a fully nonlinear model with an occasionally binding ZLB constraint (e.g., Gust et al., 2017; Plante et al., 2018; Richter and Throckmorton, 2016). This method provides the most comprehensive treatment of the ZLB constraint but is numerically intensive. It uses projection methods to solve the nonlinear model and a particle filter to evaluate the likelihood function for each draw from the posterior distribution (henceforth, NL-PF).²

The second method estimates a piecewise linear version of the nonlinear model (e.g., Guerrieri and Iacoviello, 2017). The model is solved using the OccBin toolbox developed by Guerrieri and Iacoviello (2015). The likelihood is evaluated using an inversion filter, which solves for the shocks that minimize the distance between the data and the model predictions. The benefit of this method (henceforth, OB-IF) is that it is nearly as fast as estimating a linear model with a Kalman filter while still capturing the kink in the decision rules created by the ZLB. However, OB-IF differs from NL-PF in a potentially important way. Households do not account for the possibility that the ZLB may bind in the future when it does not currently bind, which is inconsistent with survey data.³

¹Schorfheide (2000) and Otrok (2001) were the first to use these methods to generate draws from the posterior distribution of a linear DSGE model. See An and Schorfheide (2007) and Herbst and Schorfheide (2016) for examples.

²Several papers examine the effects of the ZLB constraint in a *calibrated* nonlinear model using projection methods similar to ours (e.g., Aruoba et al., 2018; Fernández-Villaverde et al., 2015; Gavin et al., 2015; Keen et al., 2017; Mertens and Ravn, 2014; Nakata, 2017; Nakov, 2008; Ngo, 2014; Richter and Throckmorton, 2015; Wolman, 2005).

³The inversion filter also removes the interest rate as an observable and sets the monetary policy shock to zero when the ZLB binds, whereas the particle filter estimates those shocks. At the ZLB, the interest rate can only identify the upper bound on policy shocks. Other observables in principle could have some additional information, but in practice there is almost none. Therefore, the particle filter does not precisely estimate these shocks when the ZLB binds.

This paper compares the accuracy of the two methods by specifying a true parameterization of a medium-scale nonlinear model with an occasionally binding ZLB constraint, solving the model with a projection method, and generating a large sample of datasets. The datasets either contain no ZLB events or a single event with various durations to understand the influence of the ZLB on our estimates. For each dataset, NL-PF and OB-IF are used to estimate a small-scale, but nested, version of the medium-scale model that generates the data. The linear model is also estimated with a Kalman filter (henceforth, Lin-KF), since that was the most common method before the Great Recession. The small-scale model excludes features of the medium-scale model that others have shown are empirically important. The difference between the two models—referred to as misspecification—accounts for the practical reality that all models are misspecified. It also sheds light on the merits of estimating a simpler, more misspecified, model with NL-PF, versus a richer, less misspecified, model with OB-IF that is numerically very costly with fully nonlinear methods.

We find NL-PF and OB-IF produce similar parameter estimates. In contrast, the predictions and forecasting performance of NL-PF are typically more accurate than OB-IF. For example, the estimates of the notional interest rate (the rate the central bank would set in the absence of the ZLB constraint), the expected ZLB duration, the probability of a 4 quarter or longer ZLB event, and forecasts of the policy rate are closer to their actual values. The increase in accuracy, however, is often small because the precautionary savings effects of the ZLB and the effects of other nonlinearities are weak in canonical models. The benefits also come with a steep increase in estimation time. The model takes roughly a week to estimate with NL-PF versus a couple hours with OB-IF.

These results suggest that OB-IF may provide an adequate substitute for NL-PF, but there are two important caveats. One, our analysis focuses exclusively on the ZLB constraint. Other constraints could create inaccuracies that provide a stronger justification for the computational burden of NL-PF. Two, OB-IF only captures nonlinearities from occasionally binding constraints. OB-IF could not account for nonlinear features such as stochastic volatility, non-convex adjustment costs, endogenous regime-switching, default, Bayesian learning, and non-Gaussian shock distributions. Our results will provide a useful benchmark for future work that examines these nonlinear features.

Model misspecification has a much larger impact on accuracy than the estimation method. It biases many of the parameter estimates and often creates significant differences between the predictions of the estimated models and the data generating process (DGP). These results suggest researchers are better off reducing misspecification by estimating a richer piecewise linear model than a simpler but computationally less intensive nonlinear model when the ZLB binds in the data. This important finding could open the door to promising new work on the implications of the ZLB.

Our paper is the first to compare different estimation methods that account for the ZLB constraint. Others compare nonlinear estimation methods to linear methods. Fernández-Villaverde and Rubio-Ramírez (2005) show that a neoclassical growth model estimated with NL-PF predicts

moments closer to the true moments than the estimates from Lin-KF using two artificial datasets and actual data. The primary source of nonlinearity in their model is high risk aversion. Hirose and Inoue (2016) generate artificial datasets from a linear model where the ZLB constraint is imposed using anticipated monetary policy shocks and then apply Lin-KF to estimate the model without the constraint. They find the estimated parameters, impulse responses, and structural shocks become less accurate as the frequency and duration of ZLB events increase in the data. Hirose and Sunakawa (2015) extend that work by generating data from a nonlinear model and re-examine the bias. None of these papers introduce misspecification, which is an important aspect of our analysis.

We also build on recent empirical work that analyzes the implications of the ZLB constraint (e.g., Gust et al., 2017; Iiboshi et al., 2018; Plante et al., 2018; Richter and Throckmorton, 2016). These papers use NL-PF to estimate a nonlinear model similar to ours using actual data from the U.S. or Japan that includes the ZLB period. Our contribution is to examine the accuracy of these nonlinear estimation methods and show under what conditions they outperform other approaches.

The measurement error (ME) in the observation equation of the filter is a key aspect of the estimation procedure that could potentially affect the accuracy of the parameter estimates. Unlike the inversion filter, the particle filter requires positive ME variances to prevent degeneracy—a situation when the likelihood is inaccurate. The literature has used a wide range of different values, with limited investigation on how they impact accuracy. Canova et al. (2014) show the downside of introducing ME is that the posterior distributions of some parameters do not contain the truth in a DSGE model estimated with Lin-KF. Cuba-Borda et al. (2017) show that ME in the particle filter reduces the accuracy of the likelihood function using a calibrated model with an occasionally binding borrowing constraint. Our analysis provides a potentially important role for ME because it includes model misspecification. We find larger ME variances improve the accuracy of some parameters, but the benefits are more than offset by decreases in the accuracy of other parameters.⁴

The paper proceeds as follows. [Section 2](#) describes the DGP and construction of our artificial datasets. [Section 3](#) outlines the estimated model and estimation methods. [Section 4](#) shows our posterior estimates and several measures of accuracy for each estimation method. [Section 5](#) concludes.

2 DATA GENERATING PROCESS

To test the accuracy of recent estimation methods that account for the ZLB constraint, we generate a large number of artificial datasets from a medium-scale New Keynesian model with capital and an occasionally binding ZLB constraint. Our model is the same as the one in Gust et al. (2017), except it removes government spending, inflation indexation, and the investment efficiency shock.⁵

⁴Herbst and Schorfheide (2018) develop a tempered particle filter that sequentially reduces the ME variances. They assess accuracy against the Kalman filter on U.S. data with a linear model and find it outperforms the untempered filter.

⁵Appendix E.7 shows how the addition of government spending to the DGP and estimated model affects our results.

2.1 FIRMS The production sector consists of a continuum of monopolistically competitive intermediate goods firms and a final goods firm. Intermediate firm $f \in [0, 1]$ produces a differentiated good, $y(f)$, according to $y_t(f) = (v_t k_{t-1}(f))^\alpha (a_t n_t(f))^{1-\alpha}$, where $n(f)$ is the labor hired by firm f and $k(f)$ is the capital rented by firm f . $a_t = z_t a_{t-1}$ is productivity and v is the capital utilization rate, which are both common across firms. Deviations from the steady-state growth rate, $\bar{z}$, follow

$$z_t = \bar{z} + \sigma_z \varepsilon_{z,t}, \quad \varepsilon_z \sim \mathbb{N}(0, 1). \quad (1)$$

The final goods firm purchases output from each intermediate firm to produce a final good, $y_t \equiv [\int_0^1 y_t(f)^{(\theta_p-1)/\theta_p} df]^{\theta_p/(\theta_p-1)}$, where $\theta_p > 1$ is the elasticity of substitution. Dividend maximization determines the demand for intermediate good f , $y_t(f) = (p_t(f)/p_t)^{-\theta_p} y_t$, where $p_t = [\int_0^1 p_t(f)^{1-\theta_p} df]^{1/(1-\theta_p)}$ is the price level. Following Rotemberg (1982), intermediate firms pay a price adjustment cost, $adj_t^p(f) \equiv \varphi_p (p_t(f)/(\bar{\pi} p_{t-1}(f)) - 1)^2 y_t/2$, where $\varphi_p > 0$ scales the cost and $\bar{\pi}$ is the steady-state gross inflation rate. Given this cost, firm f chooses $n_t(f)$, $k_{t-1}(f)$, and $p_t(f)$ to maximize the expected discounted present value of future dividends, $E_t \sum_{k=t}^{\infty} q_{t,k} d_k(f)$, subject to its production function and the demand for its product, where $q_{t,t} \equiv 1$, $q_{t,t+1} \equiv \beta(\lambda_t/\lambda_{t+1})$ is the pricing kernel between periods t and $t+1$, $q_{t,k} \equiv \prod_{j=t+1}^{k>t} q_{j-1,j}$, and $d_t(f) = p_t(f) y_t(f)/p_t - w_t n_t(f) - r_t^k v_t k_{t-1}(f) - adj_t^p(f)$. In symmetric equilibrium, the optimality conditions reduce to

$$y_t = (v_t k_{t-1})^\alpha (a_t n_t)^{1-\alpha}, \quad (2)$$

$$w_t = (1 - \alpha) m c_t y_t / n_t, \quad (3)$$

$$r_t^k = \alpha m c_t y_t / (v_t k_{t-1}), \quad (4)$$

$$\varphi_p (\pi_t / \bar{\pi} - 1) (\pi_t / \bar{\pi}) = 1 - \theta_p + \theta_p m c_t + \beta \varphi_p E_t [(\lambda_t / \lambda_{t+1}) (\pi_{t+1} / \bar{\pi} - 1) (\pi_{t+1} / \bar{\pi}) y_{t+1} / y_t], \quad (5)$$

where $\pi_t = p_t/p_{t-1}$ is the gross inflation rate. If $\varphi_p = 0$, the real marginal cost of producing a unit of output ($m c_t$) equals $(\theta_p - 1)/\theta_p$, which is the inverse of the markup of price over marginal cost.

2.2 HOUSEHOLDS Each household consists of a unit mass of members who supply differentiated types of labor, $n(\ell)$, at real wage rate $w(\ell)$. A perfectly competitive labor union bundles the labor types to produce an aggregate labor product, $n_t \equiv [\int_0^1 n_t(\ell)^{(\theta_w-1)/\theta_w} d\ell]^{\theta_w/(\theta_w-1)}$, where $\theta_w > 1$ is the elasticity of substitution. Dividend maximization determines the demand for labor type ℓ , $n_t(\ell) = (w_t(\ell)/w_t)^{-\theta_w} n_t$, where $w_t = [\int_0^1 w_t(\ell)^{1-\theta_w} d\ell]^{1/(1-\theta_w)}$ is the aggregate real wage.

The households choose $\{c_t, n_t, b_t, x_t, k_t, v_t\}_{t=0}^{\infty}$ to maximize expected lifetime utility given by $E_0 \sum_{t=0}^{\infty} \beta^t [\log(c_t - h c_{t-1}^a) - \chi \int_0^1 n_t(\ell)^{1+\eta} d\ell / (1 + \eta)]$, where β is the discount factor, χ determines steady-state labor, $1/\eta$ is the Frisch labor supply elasticity, c is consumption, c^a is aggregate consumption, h is the degree of external habit persistence, b is the real value of a privately-issued 1-period nominal bond, x is investment, and E_0 is an expectation operator conditional on information

available in period 0. Following Chugh (2006), the nominal wage rate for each labor type is subject to an adjustment cost, $adj_t^w(\ell) = \varphi_w(w_t^g(\ell) - 1)^2 y_t / 2$, where $w_t^g(\ell) = \pi_t w_t(\ell) / (\bar{\pi} \bar{z} w_{t-1}(\ell))$ is nominal wage growth relative to steady-state. The cost of utilizing the capital shock, u , is given by

$$u_t = \bar{r}^k (\exp(\sigma_v(v_t - 1)) - 1) / \sigma_v, \quad (6)$$

where $\sigma_v \geq 0$ scales the cost. Given the two costs, the household's budget constraint is given by

$$c_t + x_t + b_t / (s_t i_t) + u_t k_{t-1} + \int_0^1 adj_t^w(\ell) d\ell = \int_0^1 w_t(\ell) n_t(\ell) d\ell + r_t^k v_t k_{t-1} + b_{t-1} / \pi_t + d_t,$$

where i is the gross nominal interest rate, r^k is the capital rental rate, and d is a real dividend from ownership of intermediate firms. The nominal bond, b , is subject to a risk premium, s , that follows

$$s_t = (1 - \rho_s) \bar{s} + \rho_s s_{t-1} + \sigma_s \varepsilon_{s,t}, \quad 0 \leq \rho_s < 1, \quad \varepsilon_s \sim \mathbb{N}(0, 1), \quad (7)$$

where $\bar{s}$ is the steady-state value. An increase in s_t boosts saving, which lowers period- t demand.

Households also face an investment adjustment cost, so the law of motion for capital is given by

$$k_t = (1 - \delta) k_{t-1} + x_t (1 - \nu(x_t^g - 1)^2 / 2), \quad 0 \leq \delta \leq 1, \quad (8)$$

where $x_t^g = x_t / (\bar{z} x_{t-1})$ is investment growth relative to its steady-state and $\nu \geq 0$ scales the cost.

The first order conditions to each household's constrained optimization problem are given by

$$r_t^k = \bar{r}^k \exp(\sigma_v(v_t - 1)), \quad (9)$$

$$\lambda_t = c_t - h c_{t-1}^a, \quad (10)$$

$$w_t^f = \chi n_t^\eta \lambda_t, \quad (11)$$

$$1 = \beta E_t[(\lambda_t / \lambda_{t+1})(s_t i_t / \pi_{t+1})], \quad (12)$$

$$q_t = \beta E_t[(\lambda_t / \lambda_{t+1})(r_{t+1}^k v_{t+1} - u_{t+1} + (1 - \delta) q_{t+1})], \quad (13)$$

$$1 = q_t [1 - \nu(x_t^g - 1)^2 / 2 - \nu(x_t^g - 1)x_t^g] + \beta \nu \bar{z} E_t[(\lambda_t / \lambda_{t+1}) q_{t+1} (x_{t+1}^g)^2 (x_{t+1}^g - 1)], \quad (14)$$

$$\varphi_w(w_t^g - 1) w_t^g = [(1 - \theta_w) w_t + \theta_w w_t^f] n_t / y_t + \beta \varphi_w E_t[(\lambda_t / \lambda_{t+1}) (w_{t+1}^g - 1) w_{t+1}^g y_{t+1} / y_t], \quad (15)$$

where $1/\lambda$ is the marginal utility of consumption, q is Tobin's q , and w^f is the flexible wage rate.

Monetary Policy The central bank sets the gross nominal interest rate, i , according to

$$i_t = \max\{1, i_t^n\}, \quad (16)$$

$$i_t^n = (i_{t-1}^n)^{\rho_i} (\bar{i}(\pi_t / \bar{\pi})^{\phi_\pi} (y_t^{gdp} / (y_{t-1}^{gdp} \bar{z}))^{\phi_y})^{1-\rho_i} \exp(\sigma_i \varepsilon_{i,t}), \quad 0 \leq \rho_i < 1, \quad \varepsilon_i \sim \mathbb{N}(0, 1), \quad (17)$$

where y^{gdp} is real GDP (i.e., output, y , minus the resources lost due to adjustment costs, adj^p and

adj^w , and utilization costs), i^n is the gross notional interest rate, $\bar{i}$ and $\bar{\pi}$ are the target values of the inflation and nominal interest rates, and ϕ_π and ϕ_y are the responses to the inflation and output growth gaps. A more negative net notional rate indicates that the central bank is more constrained.

Competitive Equilibrium The aggregate resource constraint and real GDP definition are given by

$$c_t + x_t = y_t^{gdp}, \quad (18)$$

$$y_t^{gdp} = [1 - \varphi_p(\pi_t/\bar{\pi} - 1)^2/2 - \varphi_w(w_t^g - 1)^2/2]y_t - u_t k_{t-1}. \quad (19)$$

The model does not have a steady-state due to the unit root in productivity, a_t . Therefore, variables with a trend are defined in terms of productivity (i.e., $\tilde{x}_t \equiv x_t/a_t$). The detrended equilibrium system is provided in Appendix A. A competitive equilibrium consists of sequences of quantities, $\{\tilde{c}_t, \tilde{y}_t, \tilde{y}_t^{gdp}, x_t^g, y_t^g, n_t, \tilde{k}_t, \tilde{x}_t\}_{t=0}^\infty$, prices, $\{\tilde{w}_t, \tilde{w}_t^f, \tilde{w}_t^g, i_t, i_t^n, \pi_t, \tilde{\lambda}_t, v_t, u_t, q_t, r_t^k, m_{c,t}\}_{t=0}^\infty$, and exogenous variables, $\{s_t, z_t\}_{t=0}^\infty$, that satisfy the detrended equilibrium system, given the initial conditions, $\{\tilde{c}_{-1}, i_{-1}^n, \tilde{k}_{-1}, \tilde{x}_{-1}, \tilde{w}_{-1}, s_0, z_0, \varepsilon_{i,0}\}$, and three sequences of shocks, $\{\varepsilon_{z,t}, \varepsilon_{s,t}, \varepsilon_{i,t}\}_{t=1}^\infty$.

Subjective Discount Factor	β	0.9949	Rotemberg Price Adjustment Cost	φ_p	100
Frisch Labor Supply Elasticity	$1/\eta$	3	Rotemberg Wage Adjustment Cost	φ_w	100
Price Elasticity of Substitution	θ_p	6	Capital Utilization Curvature	σ_v	5
Wage Elasticity of Substitution	θ_w	6	Inflation Gap Response	ϕ_π	2
Steady-State Labor Hours	$\bar{n}$	0.3333	Output Growth Gap Response	ϕ_y	0.5
Steady-State Risk Premium	$\bar{s}$	1.0058	Habit Persistence	h	0.8
Steady-State Growth Rate	$\bar{z}$	1.0034	Risk Premium Persistence	ρ_s	0.8
Steady-State Inflation Rate	$\bar{\pi}$	1.0053	Notional Rate Persistence	ρ_i	0.8
Capital Share of Income	α	0.35	Productivity Growth Shock SD	σ_z	0.005
Capital Depreciation Rate	δ	0.025	Risk Premium Shock SD	σ_s	0.005
Investment Adjustment Cost	ν	4	Notional Interest Rate Shock SD	σ_i	0.002

Table 1: Parameter values for the data generating process.

2.3 PARAMETER VALUES Table 1 shows the model parameters for the DGP. The parameters were chosen so the DGP is characteristic of recent U.S. data. The steady-state growth rate ($\bar{z}$), inflation rate ($\bar{\pi}$), risk-premium ($\bar{s}$), and capital share of income (α) are equal to the time averages of per capita real GDP growth, the percent change in the GDP implicit price deflator, the Baa corporate bond yield relative to the yield on the 10-Year Treasury rate, and the Fernald (2012) utilization-adjusted quarterly-TFP estimates of the capital share of income from 1988Q1-2017Q4.

The subjective discount factor, β , is set to 0.9949, which is the time average of the values implied by the steady-state Euler equation and the federal funds rate. The corresponding annualized steady-state nominal interest rate is 3.3%, which is consistent with the sample average and current long-run estimates of the federal funds rate. The leisure preference parameter, χ , is set so steady-state labor equals 1/3 of the available time. The capital depreciation rate is set to 0.025. Both

values are ubiquitous in the literature. The elasticities of substitution between intermediate goods and labor types, θ_p and θ_w , are set to 6, which correspond to a 20% average markup in each sector and match the values used in Gust et al. (2017). The Frisch elasticity of labor supply, $1/\eta$, is set to 3 to match the macro estimate in Peterman (2016). The investment adjustment cost parameter, ν , and capital utilization curvature, σ_v , are consistent with the estimates in Gust et al. (2017). The price and wage adjustment cost parameters, φ_p and φ_w , are both set to 100, which correspond to Phillips curve slopes of 0.050 and 0.027. Estimates for the monetary responses to the inflation and output growth gaps, ϕ_π and ϕ_y vary in the literature, ranging from 1.5-2.5 and 0-1 (Aruoba et al., 2018; Gust et al., 2017). In this model, $\phi_\pi = 2.0$ and $\phi_y = 0.5$, which are in the middle of those ranges.

The persistence parameters and shock standard deviations are set to values that are in line with the estimates from Aruoba et al. (2018) and Gust et al. (2017). The most consequential parameters are the risk premium persistence and shock standard deviation because they have the largest impact on the expected frequency and duration of ZLB events. When either of those parameters increase, households place more weight on outcomes where the central bank cannot respond to adverse shocks by lowering the nominal interest rate, which increases the downward bias from the ZLB.

2.4 SOLUTION AND SIMULATION METHODS The nonlinear model is solved with the policy function iteration algorithm described in Richter et al. (2014), which is based on the theoretical work on monotone operators in Coleman (1991). We discretize the endogenous state variables and approximate the exogenous states, s_t , z_t , and $\varepsilon_{i,t}$ using the N -state Markov chain in Rouwenhorst (1995). The Rouwenhorst method is attractive because it only requires us to interpolate along the dimensions of the endogenous state variables, which makes the solution more accurate and faster than quadrature methods. To obtain initial conjectures for the nonlinear policy functions, we solve the level-linear analogue of our nonlinear model with Sims’s (2002) gensys algorithm. The algorithm minimizes the Euler equation errors on every node in the state space and computes the maximum distance between the updated policy functions and the initial conjectures. It then replaces the initial conjectures with the updated policy functions and iterates until the maximum distance is below the tolerance level. See Appendix B for a more detailed description of the solution method.

Data for output growth, the inflation rate, and the nominal interest rate is generated by simulating the model using the nonlinear policy functions, so the observables are given by $\mathbf{x}_t = [y_t^g, \pi_t, i_t]$. Each simulation is initialized with a draw from the ergodic distribution and contains 120 quarters, similar to what is often used when estimating models with actual data. Draws from the DGP either contain no ZLB events or a single ZLB event that is 5%, 10%, 15%, 20%, and 25% of the sample. The sample is 120 quarters, so the ZLB events are either 6, 12, 18, 24, or 30 quarters long. The longest events reflect the recent experiences of some advanced economies, such as the U.S. and Japan. There are 50 datasets for each ZLB duration. Appendix E.6 provides more information.

3 ESTIMATION METHODS

The medium-scale model is costly to estimate with global methods, which causes researchers to work with smaller models. To account for this reality, we simulate data from the fully nonlinear model and test the accuracy of different estimation methods on a small-scale nonlinear model that does not include capital or sticky wages. Therefore, the estimated model contains misspecification. The medium-scale model that generates our data collapses to the small-scale model when $\alpha = \varphi_w = 0$ and $\theta_w \rightarrow \infty$. The equilibrium system includes (1), (5), (7), (10), (12), (16), (17), and

$$y_t = a_t n_t, \quad (20)$$

$$w_t = m c_t y_t / n_t, \quad (21)$$

$$w_t = \chi n_t^\eta \lambda_t, \quad (22)$$

$$c_t = y_t^{gdp}, \quad (23)$$

$$y_t^{gdp} = [1 - \varphi_p(\pi_t/\bar{\pi} - 1)^2/2] y_t. \quad (24)$$

Once again, the trend in productivity is removed and the detrended equilibrium system is shown in Appendix A. The competitive equilibrium includes sequences of quantities, $\{\tilde{c}_t, \tilde{y}_t, \tilde{y}_t^{gdp}, y_t^g, n_t\}_{t=0}^\infty$, prices, $\{\tilde{w}_t, i_t, i_t^n, \pi_t, \tilde{\lambda}_t, m c_t\}_{t=0}^\infty$, and exogenous variables, $\{s_t, z_t\}_{t=0}^\infty$, that satisfy the detrended system, given the initial conditions, $\{\tilde{c}_{-1}, i_{-1}^n, s_0, z_0, \varepsilon_{i,0}\}$, and shock sequences, $\{\varepsilon_{z,t}, \varepsilon_{s,t}, \varepsilon_{i,t}\}_{t=1}^\infty$.

The small-scale model is estimated with Bayesian methods. For each dataset, parameters are drawn from a proposal distribution, the model is solved conditional on the draw, and a filter is applied to evaluate the likelihood function within a random walk Metropolis-Hastings algorithm. This method is used to examine the accuracy of two estimation methods that account for the ZLB.

The first method estimates the fully nonlinear model with a particle filter (NL-PF). The model is solved with the same algorithm used to generate our datasets. The filter applies Algorithm 14 in Herbst and Schorfheide (2016) by adapting the basic bootstrap particle filter described in Fernández-Villaverde and Rubio-Ramírez (2007) to include the information contained in the current observation. This allows the model to better match extreme outliers in the data. NL-PF is well-equipped to handle the nonlinearities in the data, but it is also the most computationally intensive. NL-PF requires solving the fully nonlinear model and performing a large number of simulations to evaluate the likelihood function for each draw in the random walk Metropolis-Hastings algorithm. Appendix C provides a more detailed description of the estimation algorithm and the particle filter.

The second method estimates a piecewise linear version of the nonlinear model with an inversion filter. The model is solved using the OccBin toolbox developed by Guerrieri and Iacoviello (2015). The algorithm separates the model into two regimes. In one regime, the ZLB constraint is slack, and the decision rules from the unconstrained linear model are used. In the other regime,

the ZLB binds and backwards induction within a guess and verify method solves for the decision rules. For example, if the ZLB binds in the current period, an initial conjecture is made for how many quarters the nominal interest rate will remain at the ZLB. Starting far enough in the future, the algorithm uses the decision rules for when the ZLB does not bind and iterates backward to the current period. The algorithm switches to the decision rules for the ZLB regime when the simulated nominal interest rate indicates that the ZLB binds. The simulation implies a new guess for the ZLB duration. The algorithm iterates until the implied ZLB duration equals the previous guess.

The advantage of using the piecewise linear model is that it solves very quickly. On average, the nonlinear model takes 3.6 seconds to solve (parallelized in Fortran with 16 cores), whereas the piecewise linear model takes a fraction of a second. Furthermore, the nonlinear solution time exponentially increases with the size of the model, whereas the model has little effect on the solution time in the piecewise linear model. However, it is numerically too costly to apply a particle filter. For each particle, the OccBin solution requires a long enough simulation to return to the regime where the ZLB does not bind, whereas only a 1-period update is needed with the nonlinear solution. To speed up the filter, Guerrieri and Iacoviello (2017) follow Fair and Taylor (1983) and use an inversion filter that requires only one simulation. The inversion filter solves for the shocks that minimize the distance between the observables and the equivalent model predictions each period.

The piecewise linear model estimated with the inversion filter (OB-IF) makes one potentially important simplifying assumption. Households do not account for the possibility that the ZLB may bind in the future when it does not currently bind. That means households ignore the effects of the ZLB in states of the economy where it is likely to bind in the near future because the algorithm uses the unconstrained linear decision rules. The question is whether this simplification creates large enough differences between the two methods to justify the higher estimation time of NL-PF.

As a benchmark, we estimate the linear analogue of the nonlinear model using Sims's (2002) gensys algorithm to solve the model and a Kalman filter to evaluate the likelihood function (Lin-KF). Unlike the other two methods, this method ignores the ZLB constraint, but it is much easier to implement and was the most common method used in the literature before the Great Recession.

For each estimation method, the observation equation is given by $\mathbf{x}_t = H\mathbf{s}_t + \xi_t$, where $\mathbf{s}_t$ is a vector of variables, H is an observable selection matrix, and ξ is a vector of measurement errors (MEs). The inversion filter solves for the shocks that minimize the distance between the observables, $\mathbf{x}_t$, and their model predictions, $H\mathbf{s}_t$, so there is no ME up to a numerical tolerance. With a Kalman filter or particle filter, $\xi \sim \mathcal{N}(0, R)$, where R is a diagonal matrix of ME variances.⁶ It is

⁶Ireland (2004) allows for correlated MEs, but he finds a real business cycle model's out-of-sample forecasts improve when the ME covariance matrix is diagonal. Guerrón-Quintana (2010) finds that introducing *i.i.d.* MEs and fixing the variances to 10% or 20% of the standard deviation of the data improves the empirical fit and forecasting properties of a New Keynesian model. Fernández-Villaverde and Rubio-Ramírez (2007) estimate the ME variances, but Doh (2011) argues that approach can lead to complications because the ME variances are similar to bandwidths in nonparametric estimation. Given those findings, we use a diagonal ME covariance matrix and fix the ME variances.

possible to set the ME variances to zero with the Kalman filter, since the number of observables is equal to the number of shocks. The particle filter, however, always requires positive ME variances to avoid degeneracy. Unfortunately, there is no consensus on how to set these values, despite their potentially large effect. We consider three values for the ME variances: 2%, 5%, and 10% of the variance in the data. These values encompass the wide range of values used in the literature.⁷

Parameter	Dist.	Mean (SD)	Parameter	Dist.	Mean (SD)	Parameter	Dist.	Mean (SD)
φ_p	Norm	100 (25)	h	Beta	0.8 (0.1)	σ_z	IGam	0.005 (0.005)
ϕ_π	Norm	2.0 (0.25)	ρ_s	Beta	0.8 (0.1)	σ_s	IGam	0.005 (0.005)
ϕ_y	Norm	0.5 (0.25)	ρ_i	Beta	0.8 (0.1)	σ_i	IGam	0.002 (0.002)

Table 2: Prior distributions, means, and standard deviations of the estimated parameters.

Table 2 displays information about the prior distributions of the estimated parameters. All other parameter values are fixed at their true values. The prior means are set to the true parameter values to isolate the influence of other aspects of the estimation procedure, such as the solution method and filter. Different prior means would most likely affect the accuracy of the estimation and contaminate our results. The prior standard deviations, which are consistent with the values in the literature, are relatively diffuse to give the algorithm flexibility to search the parameter space.

The estimation procedure has three stages. First, it uses a mode search to create an initial variance-covariance matrix for the estimated parameters. The covariance matrix is based on the parameters corresponding to the 90th percentile of the likelihoods from 5,000 draws. Second, it performs an initial run of the Metropolis-Hastings algorithm with 25,000 draws from the posterior distribution. The first 5,000 draws are burned off and the remaining draws are used to update the variance-covariance matrix from the mode search. Third, it conducts a final run of the Metropolis-Hastings algorithm with 50,000 draws. Our results are based on the mean draw from each dataset.

The algorithm is programmed in Fortran and the datasets are run in parallel across several supercomputers. Each dataset uses one core with OB-IF and Lin-KF, whereas NL-PF uses 16 cores because the solution is parallelized. For example, a supercomputer with 80 cores can simultaneously run 80 datasets with OB-IF but only 5 datasets with NL-PF. To increase the accuracy of the particle filter, the likelihood function is evaluated on each core. Since NL-PF uses 16 cores, we obtain 16 likelihoods and determine whether to accept a draw based on the median likelihood. This key step reduces the variance of the likelihoods from seed effects. The filter uses 40,000 particles.

⁷Some papers set the ME *standard deviations* to 20% or 25% of the sample standard deviations, which is equivalent to setting the ME *variances* to 4% or 6.25% of the sample variances (e.g., An and Schorfheide, 2007; Doh, 2011; Herbst and Schorfheide, 2016; van Binsbergen et al., 2012). Other work directly sets the ME variances to 10% or 25% of the sample variances (e.g., Bocola, 2016; Gust et al., 2017; Plante et al., 2018; Richter and Throckmorton, 2016).

	NL-PF (16 Cores)		OB-IF (1 Core)		Lin-KF (1 Core)	
	0Q	30Q	0Q	30Q	0Q	30Q
Seconds per draw	6.7 (6.1, 7.9)	8.4 (7.5, 9.5)	0.035 (0.031, 0.040)	0.096 (0.051, 0.135)	0.002 (0.002, 0.004)	0.002 (0.001, 0.003)
Hours per dataset	148.8 (134.9, 176.5)	186.4 (167.6, 210.7)	0.781 (0.689, 0.889)	2.137 (1.133, 3.000)	0.052 (0.044, 0.089)	0.049 (0.022, 0.067)

Table 3: Average and (5, 95) percentiles of the estimation times by method and ZLB duration in the data.

Table 3 shows the computing times for each estimation method. The first row reports the average and (5, 95) percentiles of the combined solution and filter times across the 50 posterior mean estimates. These draws are independent and representative of other draws. The second row shows hours per dataset, which are extrapolated by multiplying seconds per draw by 80,000 draws and dividing by 3,600 seconds per hour. Each row provides the times for NL-PF, OB-IF, and Lin-KF in datasets where the ZLB never binds and datasets with one 30 quarter ZLB event. NL-PF is run on 16 cores and the other methods use a single core. The estimation times depend on the hardware, but there are two interesting takeaways. One, OB-IF is slightly slower than Lin-KF, but it only takes a few hours to run on a single core. Two, NL-PF requires significantly more time than OB-IF, but it ran in about a week with 16 cores, so it is possible to estimate the nonlinear model on a workstation.

4 POSTERIOR ESTIMATES AND ACCURACY

The section begins by showing the accuracy of the parameter estimates for each estimation method. It then compares the filtered estimates of the notional interest rate, expected frequency and duration of the ZLB, responses to a severe recession, and the forecasting performance across the methods.

4.1 PARAMETER ESTIMATES We measure parameter accuracy by calculating the normalized root-mean square-error (NRMSE) for each estimated parameter. For parameter j and estimation method h , the error is the difference between the parameter estimate for dataset k , $\hat{\theta}_{j,h,k}$, and the true parameter, $\tilde{\theta}_j$. Therefore, the NRMSE for parameter j and estimation method h is given by

$$\text{NRMSE}_h^j = \frac{1}{\tilde{\theta}_j} \sqrt{\frac{1}{N} \sum_{k=1}^N (\hat{\theta}_{j,h,k} - \tilde{\theta}_j)^2},$$

where N is the number of datasets. The RMSE is normalized by $\tilde{\theta}_j$ to remove differences in the scales of the parameters and measure the total error. We also compute the coverage ratio given by

$$\text{CR}_h^j = \frac{1}{N} \sum_{k=1}^N \mathbb{I}(\hat{\theta}_{j,h,k}^{5\%} < \tilde{\theta}_j < \hat{\theta}_{j,h,k}^{95\%}),$$

where $\mathbb{I}$ is an indicator function and $\hat{\theta}^{X\%}$ denotes the X th percentile of the posterior distribution. This statistic shows how likely it is for the posterior distribution to contain the true parameter value.

Ptr	Truth	NL-PF-5%		OB-IF-0%		Lin-KF-5%	
		0Q	30Q	0Q	30Q	0Q	30Q
φ_p	100	151.1 (134.2, 165.8) {0.52, 0.02}	188.4 (174.7, 202.7) {0.89, 0.00}	142.6 (121.1, 157.3) {0.44, 0.08}	183.4 (169.2, 198.5) {0.84, 0.00}	151.4 (134.0, 165.7) {0.52, 0.00}	191.6 (175.3, 204.1) {0.92, 0.00}
h	0.8	0.66 (0.62, 0.70) {0.18, 0.00}	0.68 (0.64, 0.71) {0.16, 0.00}	0.64 (0.61, 0.67) {0.20, 0.00}	0.63 (0.60, 0.67) {0.21, 0.00}	0.66 (0.62, 0.69) {0.18, 0.00}	0.67 (0.63, 0.70) {0.17, 0.00}
ρ_s	0.8	0.76 (0.72, 0.80) {0.06, 0.70}	0.81 (0.78, 0.84) {0.03, 0.90}	0.76 (0.73, 0.81) {0.05, 0.82}	0.82 (0.79, 0.86) {0.04, 0.78}	0.76 (0.72, 0.80) {0.06, 0.74}	0.82 (0.78, 0.86) {0.04, 0.78}
ρ_i	0.8	0.79 (0.75, 0.82) {0.03, 0.96}	0.80 (0.75, 0.84) {0.03, 0.96}	0.76 (0.71, 0.79) {0.06, 0.52}	0.77 (0.73, 0.81) {0.05, 0.66}	0.79 (0.75, 0.82) {0.03, 0.98}	0.84 (0.80, 0.88) {0.06, 0.56}
σ_z	0.005	0.0032 (0.0023, 0.0039) {0.37, 0.00}	0.0040 (0.0030, 0.0052) {0.23, 0.58}	0.0051 (0.0044, 0.0058) {0.09, 0.92}	0.0059 (0.0050, 0.0069) {0.22, 0.30}	0.0032 (0.0023, 0.0039) {0.36, 0.00}	0.0043 (0.0030, 0.0057) {0.20, 0.68}
σ_s	0.005	0.0052 (0.0040, 0.0066) {0.15, 0.92}	0.0050 (0.0039, 0.0062) {0.13, 0.96}	0.0051 (0.0042, 0.0063) {0.13, 0.92}	0.0046 (0.0036, 0.0056) {0.15, 0.82}	0.0053 (0.0040, 0.0067) {0.15, 0.92}	0.0047 (0.0037, 0.0061) {0.15, 0.92}
σ_i	0.002	0.0017 (0.0014, 0.0020) {0.17, 0.48}	0.0015 (0.0013, 0.0019) {0.24, 0.20}	0.0020 (0.0018, 0.0023) {0.08, 0.90}	0.0020 (0.0019, 0.0024) {0.09, 0.90}	0.0017 (0.0015, 0.0020) {0.16, 0.50}	0.0016 (0.0014, 0.0019) {0.20, 0.28}
ϕ_π	2.0	2.04 (1.88, 2.19) {0.06, 0.98}	2.13 (1.94, 2.31) {0.09, 0.92}	2.01 (1.84, 2.16) {0.06, 0.98}	1.96 (1.77, 2.14) {0.06, 0.98}	2.04 (1.88, 2.20) {0.06, 0.98}	1.73 (1.52, 1.91) {0.15, 0.78}
ϕ_y	0.5	0.35 (0.21, 0.54) {0.36, 0.80}	0.42 (0.27, 0.62) {0.28, 0.98}	0.32 (0.17, 0.48) {0.41, 0.68}	0.44 (0.27, 0.61) {0.25, 0.98}	0.35 (0.22, 0.54) {0.35, 0.80}	0.32 (0.17, 0.47) {0.40, 0.76}
Σ		1.90	2.08	1.53	1.91	1.88	2.28

 Table 4: Average, (5, 95) percentiles, and {NRMSE, CR}. Σ is the sum of the NRMSE across the parameters.

Table 4 shows the parameter estimates by specification (first column header) and the duration of the ZLB (second column header). The percentage appended to each specification header corresponds to the size of the ME variances. Each cell includes the average (first row), (5, 95) percentiles (second row), NRMSE (third row, first value), and the coverage ratio (third row, second value).⁸

Across all specifications, the Rotemberg price adjustment cost parameter (φ_p) has the highest NRMSE and it becomes less accurate when the ZLB binds in the data. The upward bias is driven by misspecification, since the small-scale model used for estimation does not include sticky-wages. In the small-scale model, the response of marginal costs to shocks is much larger than in the medium-scale model, so the estimates of φ_p are higher than the true value to flatten the Phillips curve. Another inaccuracy is a downward bias in the estimates of habit persistence (h). The response of output growth to shocks is too small due to the lack of investment in the small-scale model. Lowering h increases the response to shocks, although at the expense of lower persistence. Risk premium persistence (ρ_s) and the monetary response to the output growth gap (ϕ_y) also have a downward bias in the datasets without a ZLB event, but the CR is much higher than the near-zero

⁸For conciseness, the focus is on datasets without a ZLB event and those with a 30 quarter event, but the estimates for the datasets with intermediate ZLB durations, as well as the Lin-KF-0% estimates, are provided in Appendix E.2.

values for φ_p and h . Also, the bias of ρ_s and ϕ_y decreases using datasets with a 30 quarter event.

The NL-PF-5% estimates of the productivity growth and monetary policy shock standard deviations (σ_z and σ_i) are biased downward, while the OB-IF-0% estimates are roughly consistent with their true values. In the datasets without a ZLB event, Lin-KF-5% produces identical estimates to NL-PF-5%, suggesting the bias is due to the positive ME variances in the filter. The importance of the ME variances is likely driven by the filter ascribing large shocks to ME rather than the structural shocks, reducing their estimated volatility. However, in datasets with a 30 quarter event, NL-PF-5% is more likely to contain the true risk premium parameters (ρ_s and σ_s) than OB-IF-0%. While the average estimates are similar, the CR is 0.90 for ρ_s with NL-PF-5%, compared to 0.78 with OB-IF-0%. For σ_s the CRs are 0.96 with NL-PF-5% and 0.82 with OB-IF-0%. This is notable because these two parameters have the largest effect on the frequency and duration of ZLB events.

The bottom row of [table 4](#) shows the sum of the NRMSE across the parameters. These values provide an aggregate measure of parameter accuracy. In the datasets that are not influenced by the ZLB, OB-IF-0% is more accurate than NL-PF-5%. The results for Lin-KF-5% show the lower accuracy of NL-PF-5% is driven by positive ME variances and that the ZLB is the only important nonlinearity in the model. When the ZLB binds, it reduces the accuracy of every specification, largely due to a single parameter, φ_p . Long ZLB events have the smallest effect on the accuracy of NL-PF-5%. Datasets with a 30 quarter ZLB event reduce accuracy by 0.18 relative to datasets without a ZLB event. For comparison, the accuracy decreases by 0.38 with OB-IF-0% and by 0.30 with Lin-KF-0%. However, NL-PF-5% is less accurate than OB-IF-0% due to the positive ME variances. In other words, NL-PF-5% is the best equipped to handle ZLB events in the data, but the loss in accuracy from the positive ME variances in the particle filter may outweigh those benefits.

Misspecification The absence of sticky wages and other frictions from the data generating process are important drivers of the parameter estimates in the small-scale model. Here we explore the effect of misspecification on only the OB-IF estimates since adding sticky wages substantially increases the computational cost of NL-PF. The first two columns of [table 5](#) repeat the OB-IF-0% estimates of the small-scale model, while the middle columns show the effect of reducing misspecification on the OB-IF-0% estimates by including sticky wages.⁹ The right two columns show the OB-IF-0% estimates using the medium-scale model that generates the data, eliminating all misspecification except nonlinearities not captured by the OccBin solution. For the last two cases, we fix the parameters that are not estimated in the small-scale model to their true values.¹⁰

In datasets with a 30 quarter ZLB event, adding sticky wages reduces the sum of the NRMSE from 1.91 to 1.59. That is a clear improvement over NL-PF-5% and is driven by more accurate

⁹The equilibrium system is the same as the small-scale model, except (43) and (44) are replaced with (28), (32), (40), and a real GDP definition that accounts for sticky wages (i.e., $\tilde{y}_t^{gdp} = [1 - \varphi_p(\pi_t/\bar{\pi} - 1)^2/2 - \varphi_w(w_t^g - 1)^2/2]\tilde{y}_t$).

¹⁰Appendix E.1 further explores the estimated bias by reproducing [table 4](#) with φ_p and h fixed at their true values.

Ptr	Truth	OB-IF-0%		OB-IF-0%-Sticky Wages		OB-IF-0%-DGP	
		0Q	30Q	0Q	30Q	0Q	30Q
φ_p	100	142.6 (121.1, 157.3) {0.44, 0.08}	183.4 (169.2, 198.5) {0.84, 0.00}	100.1 (76.9, 119.6) {0.13, 1.00}	129.8 (105.5, 152.3) {0.33, 0.58}	101.4 (80.1, 120.7) {0.12, 0.98}	128.4 (109.0, 148.1) {0.31, 0.46}
h	0.8	0.64 (0.61, 0.67) {0.20, 0.00}	0.63 (0.60, 0.67) {0.21, 0.00}	0.82 (0.78, 0.86) {0.04, 0.82}	0.80 (0.77, 0.85) {0.03, 0.88}	0.81 (0.75, 0.85) {0.04, 1.00}	0.77 (0.72, 0.84) {0.06, 0.78}
ρ_s	0.8	0.76 (0.73, 0.81) {0.05, 0.82}	0.82 (0.79, 0.86) {0.04, 0.78}	0.82 (0.76, 0.86) {0.04, 0.90}	0.84 (0.80, 0.88) {0.06, 0.58}	0.80 (0.76, 0.85) {0.03, 0.96}	0.82 (0.79, 0.86) {0.04, 0.80}
ρ_i	0.8	0.76 (0.71, 0.79) {0.06, 0.52}	0.77 (0.73, 0.81) {0.05, 0.66}	0.80 (0.77, 0.83) {0.02, 0.98}	0.80 (0.77, 0.84) {0.03, 0.92}	0.79 (0.75, 0.82) {0.03, 0.98}	0.79 (0.75, 0.83) {0.03, 0.92}
σ_z	0.005	0.0051 (0.0044, 0.0058) {0.09, 0.92}	0.0059 (0.0050, 0.0069) {0.22, 0.30}	0.0038 (0.0031, 0.0044) {0.24, 0.16}	0.0047 (0.0039, 0.0055) {0.12, 0.72}	0.0047 (0.0039, 0.0054) {0.11, 0.78}	0.0055 (0.0047, 0.0066) {0.15, 0.70}
σ_s	0.005	0.0051 (0.0042, 0.0063) {0.13, 0.92}	0.0046 (0.0036, 0.0056) {0.15, 0.82}	0.0085 (0.0056, 0.0134) {0.81, 0.44}	0.0074 (0.0050, 0.0107) {0.60, 0.58}	0.0060 (0.0043, 0.0084) {0.30, 0.88}	0.0051 (0.0039, 0.0068) {0.19, 0.92}
σ_i	0.002	0.0020 (0.0018, 0.0023) {0.08, 0.90}	0.0020 (0.0019, 0.0024) {0.09, 0.90}	0.0020 (0.0018, 0.0022) {0.08, 0.84}	0.0020 (0.0018, 0.0023) {0.08, 0.92}	0.0020 (0.0018, 0.0022) {0.08, 0.92}	0.0020 (0.0018, 0.0024) {0.09, 0.88}
ϕ_π	2.0	2.01 (1.84, 2.16) {0.06, 0.98}	1.96 (1.77, 2.14) {0.06, 0.98}	1.91 (1.74, 2.04) {0.07, 1.00}	1.81 (1.63, 1.99) {0.11, 0.72}	1.92 (1.72, 2.08) {0.06, 1.00}	1.81 (1.62, 2.03) {0.11, 0.70}
ϕ_y	0.5	0.32 (0.17, 0.48) {0.41, 0.68}	0.44 (0.27, 0.61) {0.25, 0.98}	0.40 (0.24, 0.58) {0.28, 0.96}	0.50 (0.33, 0.73) {0.23, 0.98}	0.41 (0.24, 0.57) {0.26, 0.96}	0.50 (0.32, 0.74) {0.24, 0.96}
Σ		1.53	1.91	1.71	1.59	1.03	1.23

 Table 5: Average, (5, 95) percentiles, and {NRMSE, CR}. Σ is the sum of the NRMSE across the parameters.

estimates of φ_p and h that dominate the lower accuracy of σ_s . The CR for φ_p and h also significantly increases. This is consistent with the claim that the bias of φ_p and h in [table 4](#) is primarily due to the lack of sticky wages, which destabilizes marginal costs and inflation. The amplification of shocks still remains too low, now for both inflation and output, which leads to an upward bias in σ_s rather than a downward bias in h . The NRMSE for σ_s is much higher and the CR declines.

Making the estimated model consistent with the DGP improves the parameter estimates even further. The sum of the NRMSE declines from 1.59 to 1.23 when the ZLB binds for 30 quarters. The primary reason is because σ_s is closer to its true value. The NRMSE in σ_s is significantly lower and the CR is much higher. Also, all of the true parameter values are encompassed by the (5, 95) percentiles of the estimates, except the estimate of φ_p has a large upward bias in the 30 quarter datasets. This result indicates that the bias of φ_p is either driven by nonlinearities not captured by the OccBin solution or the fact that datasets with long ZLB events are fairly extreme realizations of the DGP (i.e., sample selection bias).¹¹ Overall, our results suggest it is more beneficial to reduce misspecification and estimate a richer model with OB-IF than a smaller model with

¹¹Appendix E.3 shows the parameter estimates when the small-scale model is used to generate the data and estimate. The NL-PF-5% and OB-IF-0% estimates of φ_p are both close to 110 in datasets with 30 quarter ZLB events. The upward bias suggests sample selection bias also plays an important role in the medium-scale model when the ZLB binds.

NL-PF. Nonlinear methods more accurately capture the dynamics of the ZLB, but computational limitations often require excluding important features of the model, like sticky wages and capital.

Ptr	Truth	NL-PF-2%		NL-PF-5%		NL-PF-10%	
		0Q	30Q	0Q	30Q	0Q	30Q
φ_p	100	150.2 (133.5, 165.3) {0.51, 0.02}	192.0 (176.5, 207.1) {0.93, 0.00}	151.1 (134.2, 165.8) {0.52, 0.02}	188.4 (174.7, 202.7) {0.89, 0.00}	149.5 (132.6, 163.8) {0.50, 0.02}	182.7 (168.6, 197.3) {0.83, 0.02}
h	0.8	0.66 (0.62, 0.69) {0.18, 0.00}	0.67 (0.64, 0.71) {0.17, 0.00}	0.66 (0.62, 0.70) {0.18, 0.00}	0.68 (0.64, 0.71) {0.16, 0.00}	0.66 (0.61, 0.70) {0.17, 0.00}	0.68 (0.65, 0.72) {0.15, 0.00}
ρ_s	0.8	0.76 (0.71, 0.79) {0.06, 0.60}	0.81 (0.78, 0.84) {0.03, 0.92}	0.76 (0.72, 0.80) {0.06, 0.70}	0.81 (0.78, 0.84) {0.03, 0.90}	0.76 (0.72, 0.79) {0.06, 0.76}	0.81 (0.79, 0.85) {0.03, 0.88}
ρ_i	0.8	0.77 (0.73, 0.80) {0.05, 0.76}	0.79 (0.75, 0.83) {0.03, 0.96}	0.79 (0.75, 0.82) {0.03, 0.96}	0.80 (0.75, 0.84) {0.03, 0.96}	0.80 (0.77, 0.84) {0.03, 0.96}	0.81 (0.76, 0.85) {0.03, 0.94}
σ_z	0.005	0.0038 (0.0031, 0.0043) {0.25, 0.16}	0.0043 (0.0035, 0.0052) {0.18, 0.60}	0.0032 (0.0023, 0.0039) {0.37, 0.00}	0.0040 (0.0030, 0.0052) {0.23, 0.58}	0.0027 (0.0020, 0.0035) {0.46, 0.00}	0.0038 (0.0025, 0.0050) {0.28, 0.62}
σ_s	0.005	0.0052 (0.0039, 0.0065) {0.15, 0.88}	0.0051 (0.0040, 0.0061) {0.13, 0.92}	0.0052 (0.0040, 0.0066) {0.15, 0.92}	0.0050 (0.0039, 0.0062) {0.13, 0.96}	0.0051 (0.0041, 0.0065) {0.14, 0.94}	0.0049 (0.0037, 0.0061) {0.14, 0.92}
σ_i	0.002	0.0019 (0.0017, 0.0021) {0.10, 0.70}	0.0018 (0.0016, 0.0021) {0.14, 0.62}	0.0017 (0.0014, 0.0020) {0.17, 0.48}	0.0015 (0.0013, 0.0019) {0.24, 0.20}	0.0015 (0.0012, 0.0018) {0.25, 0.28}	0.0013 (0.0011, 0.0017) {0.34, 0.12}
ϕ_π	2.0	2.01 (1.84, 2.16) {0.06, 0.98}	2.14 (1.96, 2.31) {0.09, 0.90}	2.04 (1.88, 2.19) {0.06, 0.98}	2.13 (1.94, 2.31) {0.09, 0.92}	2.06 (1.89, 2.21) {0.07, 0.98}	2.12 (1.92, 2.28) {0.08, 0.96}
ϕ_y	0.5	0.31 (0.18, 0.48) {0.42, 0.64}	0.39 (0.24, 0.60) {0.32, 0.92}	0.35 (0.21, 0.54) {0.36, 0.80}	0.42 (0.27, 0.62) {0.28, 0.98}	0.41 (0.26, 0.59) {0.27, 0.98}	0.46 (0.30, 0.66) {0.24, 1.00}
Σ		1.79	2.01	1.90	2.08	1.95	2.13

Table 6: Average, (5, 95) percentiles, and {NRMSE, CR}. Σ is the sum of the NRMSE across the parameters.

ME Variances Table 6 shows the parameter estimates and NRMSEs for NL-PF with three different ME variances: 2%, 5% (baseline), and 10%. Without model misspecification, lowering the ME variances would increase the accuracy of the parameter estimates as long as the effective sample of particles is large enough. In our setup, the presence of misspecification creates a potential tradeoff. On the one hand, lower ME variances force the model to match sharp swings in the data, which could help identify the parameters. On the other hand, higher ME variances give the model a degree of freedom to account for important differences between the estimated model and the DGP.

We find higher ME variances increase the sum of the NRMSE. In datasets with 30 quarter ZLB events, it increases from 2.01 to 2.13 when the ME variances increase from 2% to 10%. For σ_z and σ_i , higher ME variances push the estimates lower, away from their true values. Once again, this result is likely driven by the filter incorrectly ascribing movements in the data to ME rather than the structural shocks. This loss in accuracy as the ME variances increase is partially offset by the increase in the accuracy of most other parameters. Estimates of ϕ_y with all datasets and estimates of φ_p with datasets where the ZLB binds for 30 quarters improve the most. These results

show that ME variances are important for accuracy. In some cases, they may compensate for model misspecification. In our setup, however, larger ME variances have a net negative effect on accuracy.

4.2 NOTIONAL INTEREST RATE ESTIMATES We measure the accuracy of the notional rate by calculating the average RMSE across periods when the ZLB binds. For period t and estimation method h , the error is the difference between the filtered notional rate based on the parameter estimates for dataset k , $\hat{i}_{t,h,k}^n$, and the true notional rate, $\tilde{i}_t^n$. The RMSE for method h is given by

$$\text{RMSE}_h^{i^n} = \sqrt{\frac{1}{N} \frac{1}{\tau} \sum_{k=1}^N \sum_{j=t}^{t+\tau-1} (\hat{i}_{j,h,k}^n - \tilde{i}_j^n)^2},$$

where t is the first period the ZLB binds and τ is the duration of the ZLB event. There is no reason to normalize the RMSE since the units are the same across periods and we do not sum across states.

Estimates of the notional interest rate are of keen interest to policymakers for two key reasons. One, they summarize the severity of the recession and the nominal interest rate policymakers would like to set in the absence of the ZLB, which help inform decisions about implementing unconventional monetary policy. Two, estimates of the notional rate help determine how long the ZLB is expected to bind, which is necessary to issue forward guidance. The notional rate is also the only latent endogenous state variable in the model that is not directly linked to an observable.

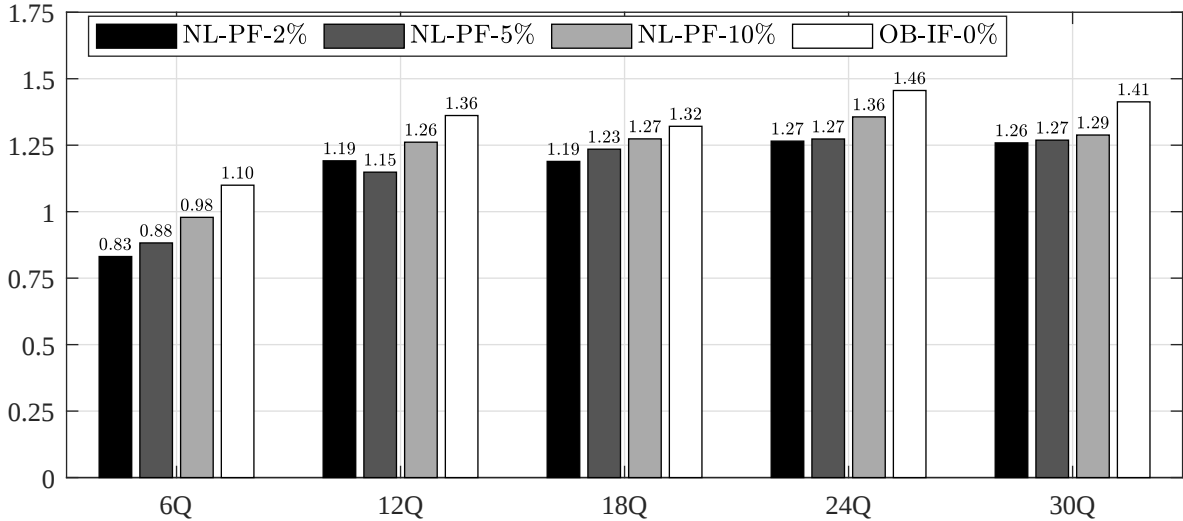


Figure 1: RMSE of the notional interest rate across ZLB durations in the data. Rates are net annualized percentages.

Figure 1 shows the accuracy of the notional rate for our baseline methods, NL-PF-5% and OB-IF-0%. It also shows how different ME variances in the particle filter affect accuracy. The Lin-KF results are not presented because they are not very informative. Since the linear model does not distinguish between the notional and nominal rates and the nominal rate is an observable, the error in the linear model equals the absolute value of the notional rate when the ZLB binds in the data.

The error in the notional rate is significant. The RMSE almost always exceeds 1 annualized percentage point and in specific periods the difference between estimated and true notional rate can exceed 2 percentage points. However, NL-PF-5% consistently provides more accurate estimates of the notional rate than OB-IF-0%. Depending on the ZLB duration, the difference between the RMSE of the two methods ranges from 0.1 to 0.25 percentage points. Appendix E.5 shows the RMSE is higher with OB-IF-0% because the estimate of the notional rate is more likely to be above the true value. It also provides an example in which the differences between the filtered notional rate paths with OB-IF-0% and NL-PF-5% reach 1 percentage point in some periods. Therefore, the average gain in accuracy from using NL-PF borders on the magnitude of policy relevance, but in some cases the differences between the two estimates can have meaningful policy implications.

4.3 EXPECTED ZLB DURATION AND PROBABILITY In addition to estimates of the notional interest rate, two commonly referenced statistics in the literature are the expected duration and probability of the ZLB constraint. These statistics determine the impact of a ZLB event in the model and are frequently measured against survey data. [Figure 2](#) shows the accuracy of the two statistics.

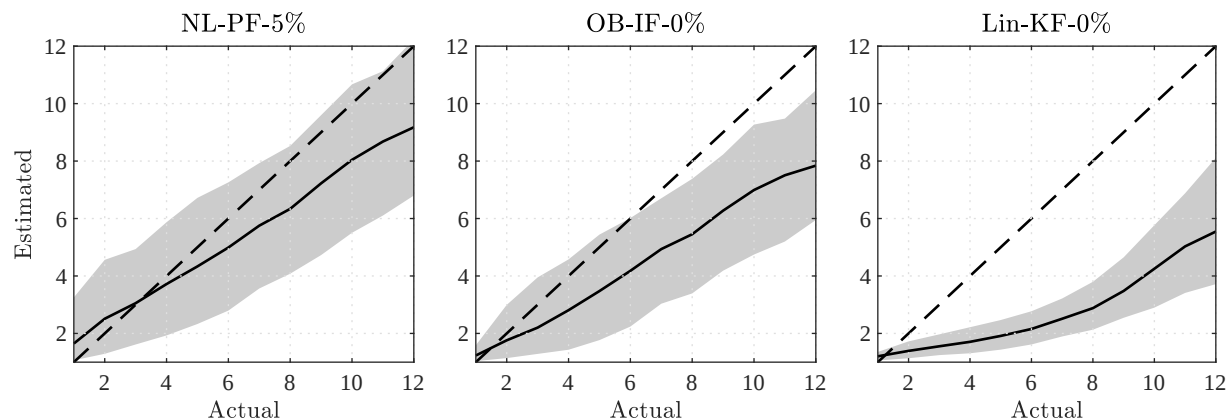
The top panel compares the expected ZLB durations given the parameter estimates from the small-scale model to the actual expected ZLB durations from the DGP given the true parameters. The expected ZLB durations are computed as the average across 10,000 simulations of a model initialized at the filtered states (or actual states for the DGP) where the ZLB binds. The solid lines are the mean expected ZLB durations in the small-scale model after pooling across the different ZLB states and datasets. The shaded areas are the (5, 95) percentiles of the durations. The estimated expected ZLB duration equals the actual expected ZLB duration along the dashed 45 degree line.

When the actual expected ZLB duration is relatively short, the NL-PF-5% and OB-IF-0% expected ZLB durations are close to the truth. As the actual expected duration lengthens, both estimates become less accurate. The NL-PF-5% 95th percentile continues to encompass the actual expected durations. However, once the actual value exceeds six quarters, there is a 95% chance or higher of under-estimating the actual expected duration with OB-IF-0%. Furthermore, the OB-IF-0% mean expected duration is typically at least one quarter shorter than the NL-PF-5% mean estimate.¹² These results are likely driven by model misspecification, as the presence of capital and sticky wages in the DGP makes the ZLB more persistent than in the estimated small-scale model.

The Lin-KF-0% estimated ZLB durations are always significantly shorter since that method does not permit a negative notional rate when filtering the data. The only instance when Lin-KF-0% produces an expected ZLB duration beyond one year is when the economy is in a severe downturn and the actual expected duration is extremely long. The Lin-KF-0% estimates are a lower bound on the OB-IF-0% estimates since the solutions are identical when the ZLB does not bind.

¹²Prior to instituting date-based forward guidance in 2011, Blue Chip consensus forecasts revealed that people expected the ZLB to bind for three quarters or less. After the forward guidance, the expectation rose to seven quarters.

(a) Estimated vs. Actual Expected ZLB Durations



(b) Estimated vs. Actual Probability of a 4 Quarter or Longer ZLB Event

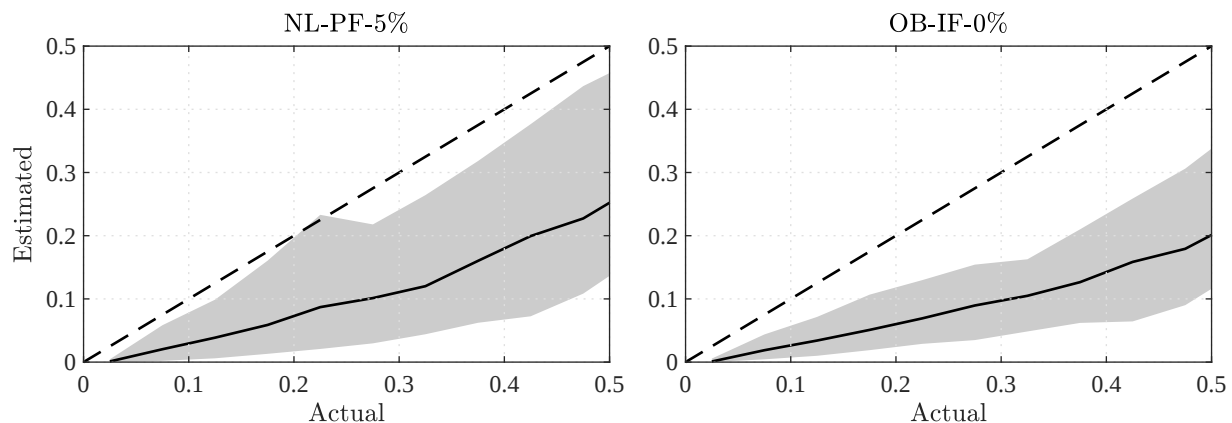


Figure 2: Estimated and actual ZLB statistics. The solid lines are mean estimates and the shaded areas capture the (5, 95) percentiles across the datasets. The dashed line shows where the estimated values would equal the actual values.

The bottom panel is constructed in a similar way as the top panel except the horizontal and vertical axes correspond to the actual and estimated probability of a ZLB event that lasts for at least four quarters. The probability is calculated in all periods where the ZLB does not bind in the data. The results for Lin-KF-0% are not shown because the probability of a four quarter ZLB event is always near zero. NL-PF-5% and OB-IF-0% underestimate the true probability, but the mean NL-PF-5% estimates are slightly closer to the actual probabilities and the 95th percentile almost encompasses the truth. Changing the ME variances in the particle filter has no discernable effect on the estimates. These results illustrate the precautionary savings effects of the ZLB, which are not captured by OB-IF-0%. However, they do not provide overwhelming support for NL-PF-5%.

4.4 RECESSION RESPONSES To illustrate the economic implications of the differences in accuracy, we compare simulations of the small-scale model given our parameter estimates to simulations of the DGP given the true parameters. The simulations are initialized in steady state and followed by four consecutive 1.5 standard deviation positive risk premium shocks, which generates

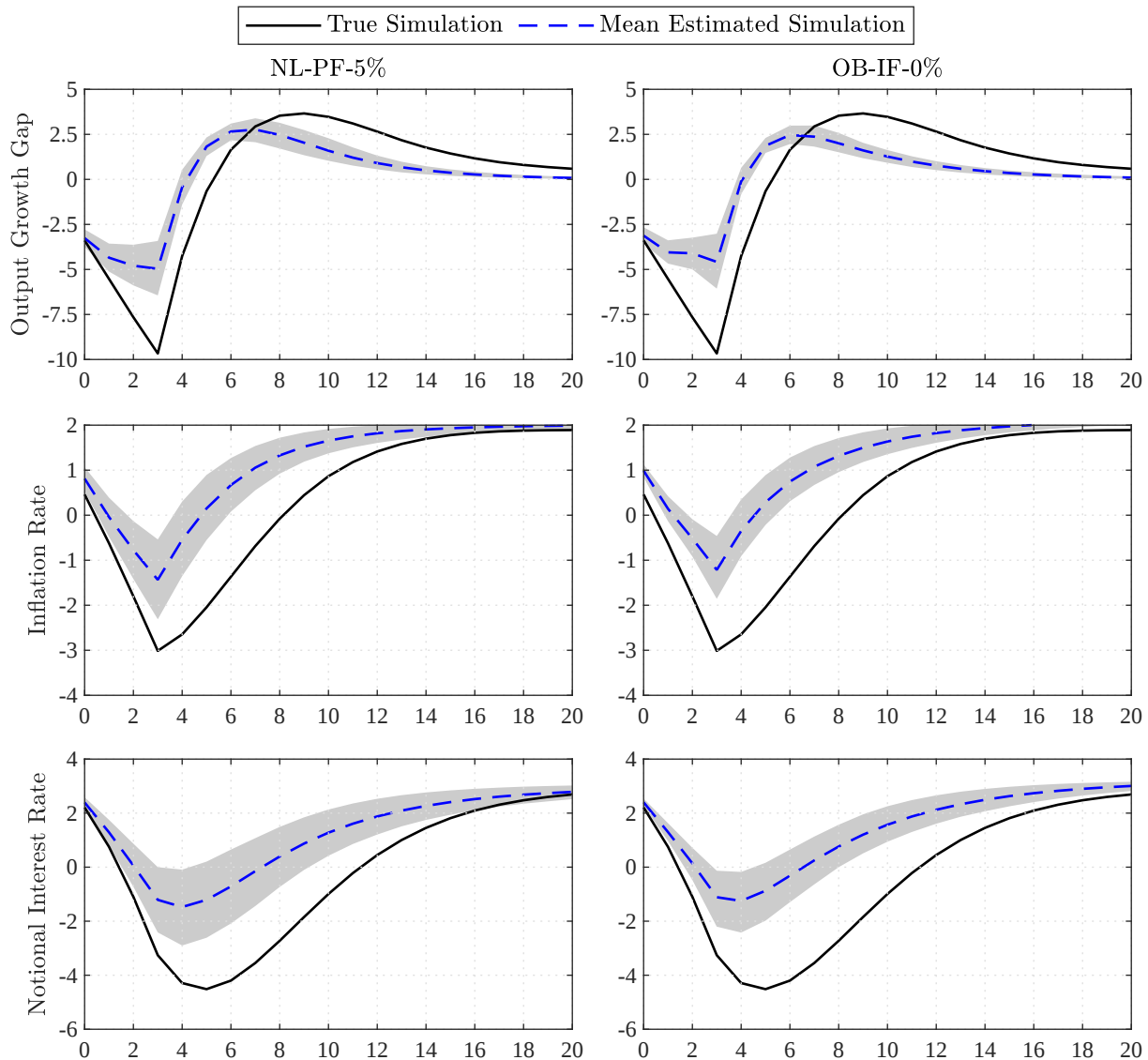


Figure 3: Recession responses. The solid line is the true simulation, the dashed line is the mean estimated simulation, and the shaded area contains the (5, 95) percentiles across the datasets. The simulations are initialized in steady state and followed by four 1.5 standard deviation positive risk premium shocks. All values are net annualized percentages.

a 10 quarter ZLB event in the DGP.¹³ A risk premium shock is a proxy for a change in demand because it affects households' consumption and saving decisions. Positive shocks cause households to postpone consumption, which reduces current output growth. We focus on this particular shock because it is the primary mechanism for generating ZLB events in the DGP and estimated model.¹⁴

Figure 3 shows the simulated paths of the output growth gap, inflation rate, and notional interest rate in annualized net percentages. The NL-PF-5% simulations are shown in the left column and

¹³The simulations are reflective of the Great Recession. The current Congressional Budget Office estimate of the output gap in 2009Q2 is -5.9% , roughly equivalent to the output (level) gap in the true simulation in the fourth period.

¹⁴Appendix E.4 shows impulse responses to a productivity growth and monetary policy shock in a severe recession.

the OB-IF-0% simulations are in the right column. The true simulation of the DGP (solid line) is compared to the mean estimated simulation of the small-scale model (dashed line). The (5, 95) percentiles account for differences in the simulations across the parameter estimates for each dataset.

Model misspecification leads to significantly muted responses relative to the true simulation.¹⁵ None of the estimated simulations for NL-PF-5% or OB-IF-0% can replicate the size of the negative output growth gap, decline in inflation, or policy response at the beginning of the true simulation. Both estimation methods also underestimate the duration of the ZLB event. However, the NL-PF-5% mean simulations of the three variables and the ZLB duration are closer to the truth than the OB-IF-0% simulations. Unlike OccBin, the fully nonlinear solution captures the expectational effects of going to the ZLB, which puts downward pressure on output and inflation and improves accuracy. Although NL-PF-5% is closer to the truth than OB-IF-0%, once again these differences are fairly small and may not justify the significantly longer estimation time.

4.5 FORECAST PERFORMANCE Another important aspect of any model is its ability to forecast. We examine the forecasting performance of each estimation method in the quarter immediately preceding a severe recession that causes the ZLB to bind. The point forecasts are inaccurate since severe recessions are rare. However, there are potentially important differences between the forecast distributions, which assign probabilities to the range of potential outcomes in a given period. The tails of the distribution are particularly important. To measure the accuracy of the forecast distribution of variable j , we compute the continuous rank probability score (CRPS) given by

$$\text{CRPS}_{m,k,t,\tau}^j = \int_{-\infty}^{\tilde{j}_{t+\tau}} [F_{m,k,t}(j_{t+\tau})]^2 dj_{t+\tau} + \int_{\tilde{j}_{t+\tau}}^{\infty} [1 - F_{m,k,t}(j_{t+\tau})]^2 dj_{t+\tau},$$

where m indicates whether the forecast distribution comes from the DGP or an estimated model, k is the dataset, t is the forecast date, $F_{m,k,t}(j_{t+\tau})$ is the cumulative distribution function (CDF) of the τ -quarter ahead forecast, and $\tilde{j}_{t+\tau}$ is the true realization. The CRPS measures the accuracy of the forecast distribution by penalizing probabilities assigned to outcomes that are not realized. It also has the same units as the forecasted variables, which are net percentages, and reduces to the mean absolute error if the forecast is deterministic. A smaller CRPS indicates a more accurate forecast.¹⁶

For each dataset, a CRPS is calculated for the small-scale model given the parameter estimates as well as the medium-scale model that generates the data. To approximate the forecast distribution for a given model, the model is simulated for 8 quarters, 10,000 times, using random shocks.¹⁷ The forecasts are initialized at the filtered state (or actual state for the DGP) one quarter before the ZLB binds in the data. The simulations are then used to approximate the CDF of the forecast distribution 8-quarters ahead. Finally, the CRPS for a given model is averaged across the datasets.

¹⁵ Appendix E.3 reproduces the responses without misspecification to confirm it is the source of the muted responses.

¹⁶ Appendix D shows the CDF for a specific dataset to illustrate what each term represents in the CRPS calculation.

¹⁷ Similar results occur with a four quarter forecast horizon, as well as with the RMSE of the point forecast.

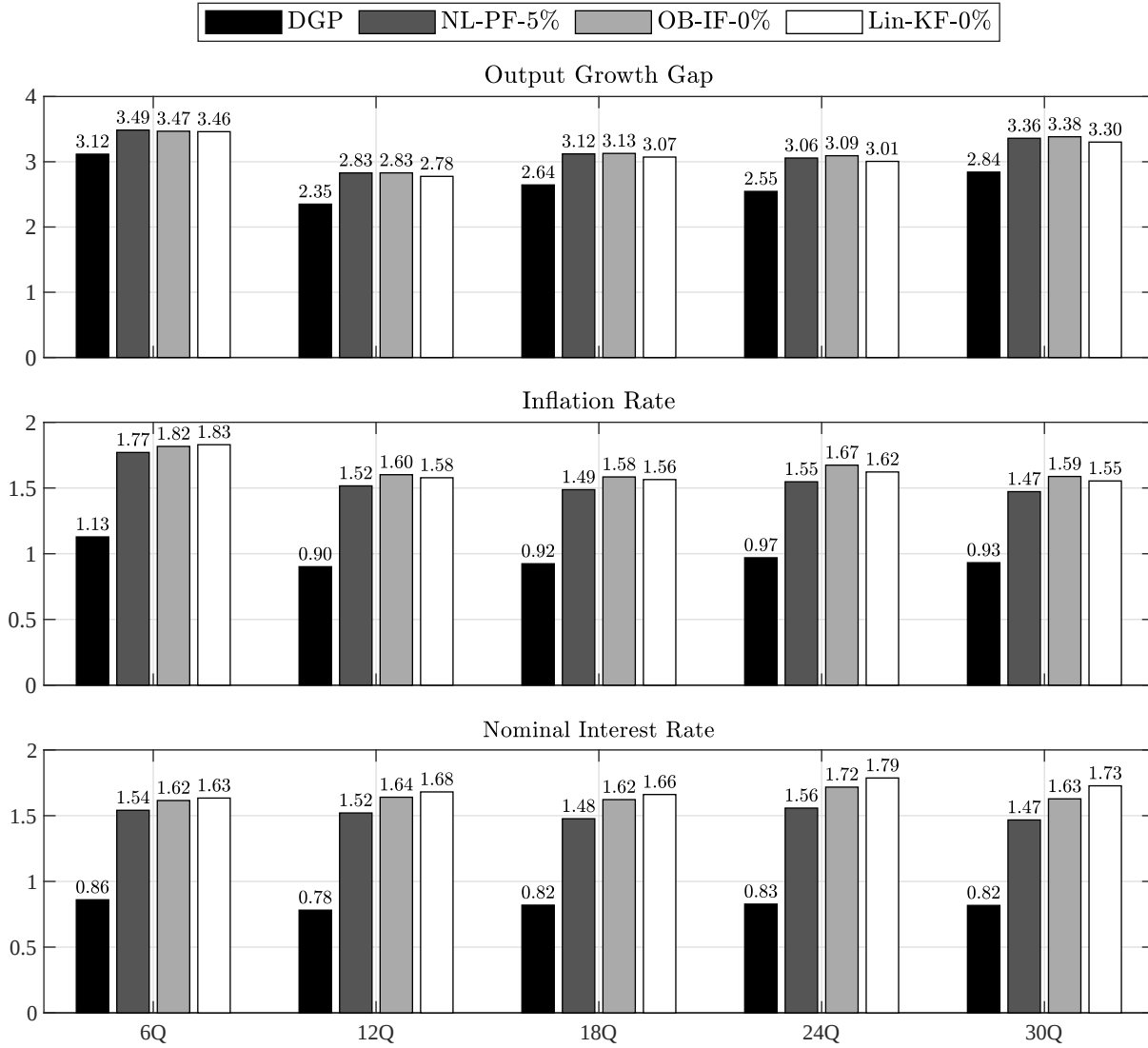


Figure 4: Mean CRPS of 8-quarter ahead forecasts. Forecasts are made one quarter before the ZLB binds in the data.

Figure 4 shows the mean CRPS across the datasets for the DGP and each estimation method. The horizontal axis denotes the ZLB duration in the data. Due to model misspecification, none of the estimation methods perform as well as the DGP. The DGP has at least a 0.5 percentage point advantage over the estimated models, regardless of the forecasted variable or ZLB duration in the data. Interestingly, the CRPS is similar across the estimation methods. The differences are most pronounced for the nominal interest rate forecasts in datasets where the ZLB binds for 30 quarters. The NL-PF-5% CRPS is only 179% of the DGP CRPS, compared to 199% for OB-IF-0% and 211% for Lin-KF-0%. The NL-PF-5% forecasts of the inflation rate are also consistently more accurate than the other estimation methods. However, in all cases the differences in accuracy are small relative to the DGP. These findings are consistent with our previous results. NL-PF-5% has an advantage over OB-IF-0%, but it is small and may not be worth the added computational costs.

5 CONCLUSION

During the Great Recession, many central banks lowered their policy rate to its ZLB, creating a kink in the policy rule and calling into question linear estimation methods. There are two promising alternatives: estimate a fully nonlinear model that accounts for the expectational effects of going to the ZLB or a piecewise linear model that is faster but ignores the expectational effects. This paper compares the accuracy of the two methods. We find the predictions of the nonlinear model are typically more accurate than the piecewise linear model, but the differences are often small. There are far larger gains in accuracy from estimating a richer, less misspecified piecewise linear model.

Our results suggest that researchers are better off using piecewise linear models rather than a simpler but properly solved nonlinear model when examining the empirical implications of the ZLB constraint. However, it is important to caution that further research is needed to examine whether our findings in the ZLB context are generalizable to other settings. It is also important to emphasize that the nonlinear model is considerably more versatile. While the piecewise linear and nonlinear models can handle any combination of occasionally binding constraints, only the nonlinear model can account for other nonlinear features emphasized in the literature (e.g., stochastic volatility, asymmetric adjustment costs, non-Gaussian shocks, search frictions, time-varying policy rules, changes in steady states). Our results will serve as an important starting point for future research that explores these nonlinear features or makes advances in nonlinear estimation methods.

REFERENCES

- AN, S. AND F. SCHORFHEIDE (2007): “Bayesian Analysis of DSGE Models,” *Econometric Reviews*, 26, 113–172, <https://doi.org/10.1080/07474930701220071>.
- ARUOBA, S., P. CUBA-BORDA, AND F. SCHORFHEIDE (2018): “Macroeconomic Dynamics Near the ZLB: A Tale of Two Countries,” *The Review of Economic Studies*, 85, 87–118, <https://doi.org/10.1093/restud/rdx027>.
- BOCOLA, L. (2016): “The Pass-Through of Sovereign Risk,” *Journal of Political Economy*, 124, 879–926, <https://doi.org/10.1086/686734>.
- CANOVA, F., F. FERRONI, AND C. MATTHES (2014): “Choosing The Variables To Estimate Singular Dsge Models,” *Journal of Applied Econometrics*, 29, 1099–1117, <https://doi.org/10.1002/jae.2414>.
- CHUGH, S. K. (2006): “Optimal Fiscal and Monetary Policy with Sticky Wages and Sticky Prices,” *Review of Economic Dynamics*, 9, 683–714, <https://doi.org/10.1016/j.red.2006.07.001>.
- COLEMAN, II, W. J. (1991): “Equilibrium in a Production Economy with an Income Tax,” *Econometrica*, 59, 1091–1104, <https://doi.org/10.2307/2938175>.

- CUBA-BORDA, P., L. GUERRIERI, M. IACOVIELLO, AND M. ZHONG (2017): “Likelihood Evaluation of Models with Occasionally Binding Constraints,” Federal Reserve Board.
- DOH, T. (2011): “Yield Curve in an Estimated Nonlinear Macro Model,” *Journal of Economic Dynamics and Control*, 35, 1229 – 1244, <https://doi.org/10.1016/j.jedc.2011.03.003>.
- FAIR, R. C. AND J. B. TAYLOR (1983): “Solution and Maximum Likelihood Estimation of Dynamic Nonlinear Rational Expectations Models,” *Econometrica*, 51, 1169–1185, <https://doi.org/10.2307/1912057>.
- FERNALD, J. G. (2012): “A Quarterly, Utilization-Adjusted Series on Total Factor Productivity,” Federal Reserve Bank of San Francisco Working Paper 2012-19.
- FERNÁNDEZ-VILLAYERDE, J., G. GORDON, P. GUERRÓN-QUINTANA, AND J. F. RUBIO-RAMÍREZ (2015): “Nonlinear Adventures at the Zero Lower Bound,” *Journal of Economic Dynamics and Control*, 57, 182–204, <https://doi.org/10.1016/j.jedc.2015.05.014>.
- FERNÁNDEZ-VILLAYERDE, J. AND J. F. RUBIO-RAMÍREZ (2005): “Estimating Dynamic Equilibrium Economies: Linear versus Nonlinear Likelihood,” *Journal of Applied Econometrics*, 20, 891–910, <https://doi.org/10.1002/jae.814>.
- (2007): “Estimating Macroeconomic Models: A Likelihood Approach,” *The Review of Economic Studies*, 74, 1059–1087, <https://doi.org/10.1111/j.1467-937X.2007.00437.x>.
- GAVIN, W. T., B. D. KEEN, A. W. RICHTER, AND N. A. THROCKMORTON (2015): “The Zero Lower Bound, the Dual Mandate, and Unconventional Dynamics,” *Journal of Economic Dynamics and Control*, 55, 14–38, <https://doi.org/10.1016/j.jedc.2015.03.007>.
- GUERRIERI, L. AND M. IACOVIELLO (2015): “OccBin: A Toolkit for Solving Dynamic Models with Occasionally Binding Constraints Easily,” *Journal of Monetary Economics*, 70, 22–38, <https://doi.org/10.1016/j.jmoneco.2014.08.005>.
- (2017): “Collateral Constraints and Macroeconomic Asymmetries,” *Journal of Monetary Economics*, 90, 28–49, <https://doi.org/10.1016/j.jmoneco.2017.06.004>.
- GUERRÓN-QUINTANA, P. A. (2010): “What You Match Does Matter: The Effects of Data on DSGE Estimation,” *Journal of Applied Econometrics*, 25, 774–804, <https://doi.org/10.1002/jae.1106>.
- GUST, C., E. HERBST, D. LÓPEZ-SALIDO, AND M. E. SMITH (2017): “The Empirical Implications of the Interest-Rate Lower Bound,” *American Economic Review*, 107, 1971–2006, <https://doi.org/10.1257/aer.20121437>.
- HERBST, E. AND F. SCHORFHEIDE (2018): “Tempered Particle Filtering,” *Journal of Econometrics*, forthcoming, <https://doi.org/10.1016/j.jeconom.2018.11.003>.
- HERBST, E. P. AND F. SCHORFHEIDE (2016): *Bayesian Estimation of DSGE Models*, Princeton, NJ: Princeton University Press.

- HIROSE, Y. AND A. INOUE (2016): “The Zero Lower Bound and Parameter Bias in an Estimated DSGE Model,” *Journal of Applied Econometrics*, 31, 630–651, <https://doi.org/10.1002/jae.2447>.
- HIROSE, Y. AND T. SUNAKAWA (2015): “Parameter Bias in an Estimated DSGE Model: Does Nonlinearity Matter?” Centre for Applied Macroeconomic Analysis Working Paper 46/2015.
- IIBOSHI, H., M. SHINTANI, AND K. UEDA (2018): “Estimating a Nonlinear New Keynesian Model with a Zero Lower Bound for Japan,” Tokyo Center for Economic Research Working Paper E-120.
- IRELAND, P. N. (2004): “A Method for Taking Models to the Data,” *Journal of Economic Dynamics and Control*, 28, 1205–1226, [https://doi.org/10.1016/S0165-1889\(03\)00080-0](https://doi.org/10.1016/S0165-1889(03)00080-0).
- KEEN, B. D., A. W. RICHTER, AND N. A. THROCKMORTON (2017): “Forward Guidance and the State of the Economy,” *Economic Inquiry*, 55, 1593–1624, <https://doi.org/10.1111/ecin.12466>.
- LAUBACH, T. AND J. C. WILLIAMS (2016): “Measuring the Natural Rate of Interest Redux,” Finance and Economics Discussion Series 2016-011.
- MERTENS, K. AND M. O. RAVN (2014): “Fiscal Policy in an Expectations Driven Liquidity Trap,” *The Review of Economic Studies*, 81, 1637–1667, <https://doi.org/10.1093/restud/rdu016>.
- NAKATA, T. (2017): “Uncertainty at the Zero Lower Bound,” *American Economic Journal: Macroeconomics*, 9, 186–221, <https://doi.org/10.1257/mac.20140253>.
- NAKOV, A. (2008): “Optimal and Simple Monetary Policy Rules with Zero Floor on the Nominal Interest Rate,” *International Journal of Central Banking*, 4, 73–127.
- NGO, P. V. (2014): “Optimal Discretionary Monetary Policy in a Micro-Founded Model with a Zero Lower Bound on Nominal Interest Rate,” *Journal of Economic Dynamics and Control*, 45, 44–65, <https://doi.org/10.1016/j.jedc.2014.05.010>.
- OTROK, C. (2001): “On Measuring the Welfare Cost of Business Cycles,” *Journal of Monetary Economics*, 47, 61–92, [https://doi.org/10.1016/S0304-3932\(00\)00052-0](https://doi.org/10.1016/S0304-3932(00)00052-0).
- PETERMAN, W. B. (2016): “Reconciling Micro and Macro Estimates of the Frisch Labor Supply Elasticity,” *Economic Inquiry*, 54, 100–120, <https://doi.org/10.1111/ecin.12252>.
- PLANTE, M., A. W. RICHTER, AND N. A. THROCKMORTON (2018): “The Zero Lower Bound and Endogenous Uncertainty,” *Economic Journal*, 128, 1730–1757, <https://doi.org/10.1111/ecoj.12445>.
- RICHTER, A. W. AND N. A. THROCKMORTON (2015): “The Zero Lower Bound: Frequency, Duration, and Numerical Convergence,” *B.E. Journal of Macroeconomics*, 15, 157–182, <https://doi.org/10.1515/bejm-2013-0185>.
- (2016): “Is Rotemberg Pricing Justified by Macro Data?” *Economics Letters*, 149, 44–48, <https://doi.org/10.1016/j.econlet.2016.10.011>.

- RICHTER, A. W., N. A. THROCKMORTON, AND T. B. WALKER (2014): “Accuracy, Speed and Robustness of Policy Function Iteration,” *Computational Economics*, 44, 445–476, <https://doi.org/10.1007/s10614-013-9399-2>.
- ROTEMBERG, J. J. (1982): “Sticky Prices in the United States,” *Journal of Political Economy*, 90, 1187–1211, <https://doi.org/10.1086/261117>.
- ROUWENHORST, K. G. (1995): “Asset Pricing Implications of Equilibrium Business Cycle Models,” in *Frontiers of Business Cycle Research*, ed. by T. F. Cooley, Princeton, NJ: Princeton University Press, 294–330.
- SCHORFHEIDE, F. (2000): “Loss Function-Based Evaluation of DSGE Models,” *Journal of Applied Econometrics*, 15, 645–670, <https://doi.org/10.1002/jae.582>.
- SIMS, C. A. (2002): “Solving Linear Rational Expectations Models,” *Computational Economics*, 20, 1–20, <https://doi.org/10.1023/A:1020517101123>.
- VAN BINSBERGEN, J. H., J. FERNÁNDEZ-VILLAYERDE, R. S. KOIJEN, AND J. RUBIO-RAMÍREZ (2012): “The Term Structure of Interest Rates in a DSGE Model with Recursive Preferences,” *Journal of Monetary Economics*, 59, 634–648, <https://doi.org/10.1016/j.jmoneco.2012.09.002>.
- WOLMAN, A. L. (2005): “Real Implications of the Zero Bound on Nominal Interest Rates,” *Journal of Money, Credit and Banking*, 37, 273–96, <https://doi.org/10.1353/mcb.2005.0026>.