

# Investment Nonlinearity in Smets–Wouters: A Switching-Fidelity Approach\*

Mátyás Farkas<sup>†</sup>

International Monetary Fund

August 3, 2026

*Preliminary — Comments Welcome*

## Abstract

Where do nonlinear dynamics matter in a medium-scale model? I compare the nonlinear Harding–Lindé–Trabandt version of Smets–Wouters with its first-order approximation at the same parameters, states, and shocks. Direct stochastic extended-path solutions show that investment accounts for 78 percent of the squared gap across observables in model units and 66 percent after scaling each series by its nonlinear variation. Parameter ablations trace the gap to investment adjustment and capital utilization; price-setting nonlinearity matters only under extreme calibrations. In a common investment-shock experiment, higher-order local solutions close at most 16 percent of the nonlinear path gap. Switching-fidelity estimation makes the direct comparison practical: a supervised neural network learns the nonlinear-minus-linear gap and reduces raw pooled transition error by 43 percent on parameter-held-out validation paths. The method makes approximation error and failed nonlinear solves visible, but it audits one local solution branch rather than proving global uniqueness.

**JEL Classification:** C11, C13, C32, C63, E47

**Keywords:** DSGE models, nonlinear estimation, investment adjustment costs, perturbation bias, switching-fidelity estimation, neural residual surrogates

---

\*I am grateful to Fabio Canova, Kai Christoffel, Chris Erceg, Jesper Lindé, Zoltán Jakab, Thore Klockerols, and Marcin Kolasa for insightful discussions and comments. All errors are my own. The views expressed herein are those of the author and not necessarily those of the International Monetary Fund.

<sup>†</sup>mfarkas@imf.org

# Contents

|          |                                                      |           |
|----------|------------------------------------------------------|-----------|
| <b>1</b> | <b>Introduction</b>                                  | <b>4</b>  |
| <b>2</b> | <b>How Existing Methods Handle the Cost</b>          | <b>6</b>  |
| 2.1      | Higher-Order and Pathwise Approximations . . . . .   | 6         |
| 2.2      | Particle and Simulation-Based Likelihoods . . . . .  | 7         |
| 2.3      | Filter-Free Joint Sampling . . . . .                 | 7         |
| 2.4      | Neural Solutions and Likelihood Surrogates . . . . . | 8         |
| 2.5      | When Switching Fidelity Is Useful . . . . .          | 9         |
| <b>3</b> | <b>What First Order Misses</b>                       | <b>9</b>  |
| 3.1      | A Simple Local Result . . . . .                      | 9         |
| 3.2      | A Local Ranking . . . . .                            | 10        |
| <b>4</b> | <b>Models and Estimation</b>                         | <b>11</b> |
| 4.1      | General Framework . . . . .                          | 11        |
| 4.2      | Observation Equation . . . . .                       | 11        |
| 4.3      | Estimation Problem . . . . .                         | 12        |
| 4.4      | The Smets–Wouters–HLT Model . . . . .                | 13        |
| <b>5</b> | <b>How Switching Fidelity Works</b>                  | <b>13</b> |
| 5.1      | Residual Correction . . . . .                        | 13        |
| 5.2      | Stochastic Extended Path (SEP) Algorithm . . . . .   | 14        |
| 5.3      | Dataset Generation and Surrogate Training . . . . .  | 14        |
| 5.4      | Local Solution Selection and Its Limit . . . . .     | 15        |
| 5.5      | Regime-Switching Gate and Inversion Filter . . . . . | 16        |
| 5.6      | Filter-Free Hamiltonian Monte Carlo . . . . .        | 16        |
| <b>6</b> | <b>Where the Nonlinearity Comes From</b>             | <b>17</b> |

|          |                                                             |           |
|----------|-------------------------------------------------------------|-----------|
| 6.1      | How I Measure the Gap . . . . .                             | 17        |
| 6.2      | Is the Lower Bound the Main Source? . . . . .               | 19        |
| 6.3      | Can Higher-Order Local Solutions Close the Gap? . . . . .   | 22        |
| 6.4      | An Investment-Sensitive Policy Experiment . . . . .         | 23        |
| 6.5      | Removing Investment Adjustment Costs . . . . .              | 23        |
| <b>7</b> | <b>What Changes in Estimation</b>                           | <b>25</b> |
| 7.1      | The Correction Relative to Parameter-Matched ROM1 . . . . . | 25        |
| 7.2      | What HMC Reveals . . . . .                                  | 26        |
| 7.3      | An Implementation Ablation Inside the HMC Target . . . . .  | 27        |
| 7.4      | Computation . . . . .                                       | 27        |
| <b>8</b> | <b>Robustness and Scope</b>                                 | <b>28</b> |
| 8.1      | Does the Result Depend on Price Setting? . . . . .          | 28        |
| 8.2      | Does Adjustment-Cost Curvature Change the Result? . . . . . | 28        |
| 8.3      | Scope of the Evidence . . . . .                             | 29        |
| 8.4      | What the Result Means . . . . .                             | 30        |
| <b>9</b> | <b>Conclusion</b>                                           | <b>31</b> |

# 1 Introduction

First-order approximations are the workhorse of estimated macroeconomic models. Their attraction is obvious: they are fast, stable, and accurate near steady state. The same approximation, however, makes large shocks look like scaled-up small shocks and positive shocks like mirror images of negative ones. Whether that loss matters depends on which part of the model it removes.

In the Smets–Wouters model, the largest gap appears where one might not expect it. The model is known for nominal rigidities and the lower bound on interest rates, yet its largest departure from first order appears in investment. I obtain this result by solving the model twice at the same state, shock, and parameter vector. The first solution is the usual first-order approximation; the second solves the nonlinear equilibrium with stochastic extended path. Among the seven series used to estimate the model, investment accounts for 78.4 percent of the squared difference in model units and 66.2 percent after each series is scaled by its variation under the nonlinear solution. Inflation accounts for less than 0.4 percent under either measure. For the variables used in estimation, the important nonlinearity is therefore real, not nominal.

The source of the result is the investment block itself. At steady growth, the level adjustment cost and its slope are both zero. A first-order expansion then loses the resource cost of changing investment, together with the interactions among marginal adjustment costs, capital values, and discounting. Small changes in investment adjustment costs or capital utilization move the nonlinear gap more than similar changes in price and wage curvature at almost every parameter vector I examine. Price curvature becomes important only under deliberately extreme calibrations. The evidence thus points to the capital side of the model, rather than to price setting.

The obstacle is computation. A direct nonlinear transition takes about 0.8 seconds, too long to repeat at every step of an estimation. Switching fidelity separates the expensive calculation from the sampler. I use direct nonlinear solutions offline to measure the error of

the first-order model, then train a neural network to learn that error. The sampler applies the correction in stressed states and keeps the first-order solution in calm ones.

The correction has a simple benchmark. Below, ROM1 denotes the ungated first-order solution and stochastic extended path (SEP) the nonlinear solution. I re-solve ROM1 at each parameter vector and hold twelve complete 18-parameter vectors out of training. On their 2,208 transitions, adding the learned residual lowers raw pooled observable RMSE from 2.19 to 1.24, a 43 percent reduction. Every observable improves. This comparison changes only the correction: the state, shock, and parameter vector are identical, and neither the gate nor the inversion procedure enters.

Supervision also fixes the numerical object that the network learns. SEP supplies one direct label at each accepted input; failed solves are mapped rather than filled in. Nearby parameter ablations at three shock sizes remain on the same convergent branch. These checks make the local solution choice explicit. They do not prove that the nonlinear model has no other branch, basin, or large-shock instability. That global question is separate from approximation accuracy.

Hamiltonian Monte Carlo adds another form of transparency. Four matched chains show the joint uncertainty, correlations, effective sample sizes, Monte Carlo error, energy behavior, and difficult directions of the 18-parameter profiled target. These diagnostics establish mixing within the initialized basin. They do not search for other economic equilibria or posterior modes.

The design follows reduced-order modeling by learning the error of a cheap solver rather than replacing the nonlinear model itself (Peherstorfer et al., 2018). Stochastic extended path supplies the direct benchmark (Fair and Taylor, 1983; Adjemian and Juillard, 2025), while filter-free sampling carries the corrected transition into estimation without a particle filter (Childers et al., 2022; Farkas and Tatár, 2020). The paper moves from the cost of nonlinear estimation to the investment mechanism, the switching-fidelity solution, the HMC evidence, and the boundary between local validation and global stability.

## 2 How Existing Methods Handle the Cost

Nonlinear estimation shifts the computational burden among model solution, integration, sampling, and approximation. No method removes the burden; each chooses where to bear it. The useful choice therefore depends on the economic question and on which approximation error can be checked directly.

### 2.1 Higher-Order and Pathwise Approximations

Second- and third-order perturbations recover risk and curvature terms that are absent at first order (Schmitt-Grohé and Uribe, 2004; Kim et al., 2008). Pruning removes the artificial explosive terms generated by repeated simulation of a polynomial decision rule and gives a stable state-space representation (Kim et al., 2008; Andreasen et al., 2018). These methods remain local: their derivatives are computed at one expansion point. Accuracy can therefore deteriorate after large shocks or far from the steady state. In a rare-disaster New Keynesian model, even third-order perturbation has sizable accuracy errors, while global Smolyak collocation is costly in time and memory (Fernández-Villaverde and Levintal, 2018).

Piecewise-linear methods handle known occasionally binding constraints by solving a sequence of regimes (Guerrieri and Iacoviello, 2015; Aruoba et al., 2021). They are well suited to kinks, but they do not recover general smooth curvature within a regime. Dynamic perturbation instead recomputes local derivatives along the equilibrium path and can handle large transitions and binding constraints (Mennuni et al., 2025). Its cost is pathwise: the model is solved along each requested path rather than once as a global policy function. Related estimation work recomputes local linear solutions around forecast states and evaluates the resulting time-varying linear system with a Kalman filter (Janssens and McCrary, 2025). This avoids a global grid and a nonlinear filter, but its accuracy remains tied to the sequence of local linearizations. Projection and collocation methods approximate a policy over a larger state region, but the grid or basis grows quickly with the number of states

(Judd, 1992; Malin et al., 2011; Fernández-Villaverde and Levintal, 2018).

I keep ROM1 as the baseline and use SEP only to measure what ROM1 misses. SEP is therefore outside the sampler. The saving is large, but the guarantee is local: the direct checks cover a finite set of states, shocks, and parameters.

## 2.2 Particle and Simulation-Based Likelihoods

Particle filters provide likelihood-based inference for nonlinear and non-Gaussian state-space models (Fernández-Villaverde and Rubio-Ramírez, 2007). In the large-particle limit they target the nonlinear filtering distribution. The practical problem is weight concentration. When the proposal poorly matches an informative observation, or when the effective dimension is high, most weight can fall on very few particles; the required ensemble size can then grow very quickly (Snyder et al., 2008). A large shock can make this mismatch worse, but collapse is not caused by large shocks alone.

Better proposals, tempering, resampling, and Markov-chain rejuvenation can substantially improve a bootstrap filter (Herbst and Schorfheide, 2019; Andrieu et al., 2010). They also add computation and design choices. Parameter estimation compounds the cost: particle MCMC embeds a state filter in each parameter proposal, while SMC<sup>2</sup> propagates parameter particles that each carry a state filter and periodically rejuvenates them (Andrieu et al., 2010; Chopin et al., 2013).

I do not use particle weights. I recover shocks from the observation equation or sample them only in stressed periods. This keeps the objective differentiable and low dimensional. It also changes the statistical object. Profiling shocks is not the same as integrating them out, so the resulting objective is not a marginal likelihood.

## 2.3 Filter-Free Joint Sampling

Filter-free HMC treats structural parameters, latent innovations or states, and initial conditions as one differentiable posterior target (Childers et al., 2022). This avoids sequential

importance weights and has worked in nonlinear DSGE applications. [Naubert \(2025\)](#), for example, uses a mixture density network for the initial-state distribution in a medium-scale nonlinear model.

Observation inversion offers a lower-dimensional special case. When the number of observables equals the number of shocks and the observation equation can be inverted recursively, shocks can be recovered instead of integrated out ([Kollmann, 2017](#)). Multiple roots, measurement error, or a mismatch between observables and shocks require an additional rule or approximation.

The cost moves from filtering to dimension. Sampling  $K$  shocks over  $T$  dates adds  $TK$  unknowns, and every HMC gradient passes through the full state path. I reduce that cost in two ways. The gated version samples shocks only in stressed periods. The inversion version profiles them out. A bad gate can still miss an important shock, and profiling still replaces integration with optimization.

## 2.4 Neural Solutions and Likelihood Surrogates

Neural networks can approximate policy functions in state spaces that are too large for conventional grids ([Maliar et al., 2021](#); [Azinovic et al., 2022](#)). In unsupervised equilibrium networks, training points come from simulated paths and the loss is built from equilibrium conditions. This makes labels cheap, but the researcher must choose the state sampler, loss weights, architecture, optimizer, and stopping rule. Small equilibrium residuals alone need not prove that the approximation is close to the intended equilibrium ([Kubler, 2011](#)). The concern is solution selection and verification, not a claim that unsupervised networks generally fail.

Recent work also learns the nonlinear policy and a likelihood surrogate. [Kase et al. \(2025\)](#) train neural policy functions for a nonlinear HANK model and emulate particle-filter likelihood evaluations over parameter draws. This can amortize a demanding global solution and filtering exercise, but the accuracy of both networks must be checked over the parameter

and state region used in estimation.

I instead train on a measured residual. SEP supplies the label  $g^{\text{SEP}} - g^{\text{ROM1}}$ , ROM1 supplies the fallback, and held-out SEP solves reveal the error. New labels are added where the sampler travels. An unsupervised equilibrium network may recover the same branch, but its residual loss does not by itself provide an external branch label or a map of failed direct solves. Supervision makes both objects observable before estimation. It does not remove the training choices, the cost of SEP labels, or the need to define a finite support region.

## 2.5 When Switching Fidelity Is Useful

Switching fidelity is useful when ROM1 works well in most periods and its error is concentrated in a smaller region that SEP can cover. It is a poor choice when there is no reliable baseline, when SEP fails in the relevant region, or when the research question requires an exact marginal likelihood. The method saves online computation by accepting a finite, explicit accuracy claim.

## 3 What First Order Misses

The numerical ranking has a simple analytical basis. If a nonlinear term and its slope are both zero at steady state, a first-order solution cannot see that channel.

### 3.1 A Simple Local Result

**Proposition 3.1** (Perturbation Invisibility). *Let a smooth equilibrium system be written as  $F(z_t, f(x_t); \theta) = 0$ , with steady state  $(\bar{z}, \bar{x})$ . Assume that the Jacobian of  $F$  with respect to the endogenous variables  $z_t$  is nonsingular at the steady state, and that  $f$  enters through the displayed differentiable argument. If  $f(\bar{x}) = 0$  and  $f'(\bar{x}) = 0$ , then this channel adds no term*

to the first-order perturbation coefficients. Specifically, with  $\hat{x}_t = x_t - \bar{x}$ ,

$$f(x_t) = \frac{1}{2}f''(\bar{x})\hat{x}_t^2 + \mathcal{O}(\hat{x}_t^3),$$

so the derivative of the composite equilibrium condition through  $f$  is zero at first order.

The proposition applies to one channel at a time. A function may still matter at first order if it enters another equation through a nonzero derivative. The point is narrower: a term with zero level and zero slope enters this equation only at second order.

Investment adjustment costs meet both conditions. Write  $S(x) = \frac{\chi}{2}(x - 1)^2$ , where  $x = I_t/I_{t-1}$ . At trend growth,  $S(1) = S_x(1) = 0$ . First order therefore drops the level resource cost. It retains the local slope of the marginal cost, but not its products with Tobin's  $q$ , the capital price, and the discount factor.

Capital utilization is different. Its cost satisfies  $a(1) = 0$  but  $a'(1) = r_{ss}^k \neq 0$ . First order captures the main utilization response and misses only curvature. That curvature is about 1,200 times smaller than  $S''(1)$  at the baseline calibration.

Kimball aggregation shows why curvature alone is not enough. Prices and wages have high formal curvature, but their dispersion indices stay within  $10^{-4}$  of one under the baseline Calvo calibration. A curved function matters only if its argument moves far enough from steady state.

## 3.2 A Local Ranking

I rank seven nonlinear terms by the size of their quadratic term relative to their linear term at a one-standard-deviation displacement. The investment first-order condition ranks first at 18.7 percent. Its local slope is visible at first order, but its interactions with capital prices and discount factors are not.

The next two terms are much smaller. Wage setting reaches 2.8 percent because wages are sticky and wage dispersion moves more than price dispersion. Habit in consumption

reaches 2.7 percent because habit makes the relevant consumption base small. All other terms are at or below 1.2 percent. In the HLT simulations below, the lower bound affects the gap when it binds, but binding periods are too rare to drive the average result.

The ranking predicts that the investment block should dominate the SEP–ROM1 gap, the wage block should contribute modestly, and the pricing block and ELB should be small at empirically relevant shock scales. Section 6 finds patterns consistent with these predictions.

## 4 Models and Estimation

### 4.1 General Framework

Let  $s_t$  collect the model’s states and jump variables. The nonlinear equilibrium conditions are

$$\mathbb{E}_t [F(s_{t+1}, s_t, s_{t-1}, \varepsilon_t; \theta)] = 0, \quad (1)$$

where  $s_t$  collects states and jumps,  $\varepsilon_t$  are standardized shocks, and  $\theta$  are structural parameters. A solution is a transition rule

$$s_t = g_\theta(s_{t-1}, \varepsilon_t),$$

that satisfies (1). The paper compares two transition rules: the first-order perturbation  $g_\theta^{\text{ROM1}}$ , which I call ROM1, and the direct nonlinear stochastic-extended-path solution  $g_\theta^{\text{SEP}}$ . I use full-order model (FOM) as shorthand for this direct SEP benchmark.

### 4.2 Observation Equation

Observed data are

$$y_t = H(s_t) + \eta_t, \quad \eta_t \sim \mathcal{N}(0, \Sigma_y), \quad (2)$$

with a linear measurement map  $H$ . In the Smets–Wouters application, the observables are output growth, consumption growth, investment growth, hours, inflation, wage growth, and the interest rate.

### 4.3 Estimation Problem

For reference, an exact Bayesian analysis targets

$$p(\theta \mid y_{1:T}^{\text{obs}}) \propto p(y_{1:T}^{\text{obs}} \mid \theta) p(\theta), \quad (3)$$

with the ROM1/Kalman likelihood providing an exact linear-Gaussian instance. In small validation models I sample shocks jointly with  $\theta$ . The full latent-shock density would be

$$\mathcal{L}(\theta, s_0, \varepsilon_{1:T} \mid y_{1:T}^{\text{obs}}) = \prod_{t=1}^T p(y_t^{\text{obs}} \mid s_t; \Sigma_y) \cdot p(s_0) \cdot \prod_{t=1}^T p(\varepsilon_t),$$

where  $s_t = \widehat{g}_{\theta,w}(s_{t-1}, \varepsilon_t)$ .

The full 18-parameter exercise uses a different object. In calm quarters it uses the Kalman contribution. In stressed quarters it chooses the shock that best fits the data and evaluates the observation error at that shock. Let  $m_t$  be the fixed gate,  $\widehat{\varepsilon}_t(\theta)$  the recovered shock, and  $b \in \{\text{NN}, 0\}$  indicate whether the neural residual is active. The objective is

$$Q_{T,b}^{\text{prof}}(\theta) = \sum_{t:m_t=0} \ell_t^K(\theta) + \sum_{t:m_t=1} \ell_{t,b}^{\text{obs}}(\theta; \widehat{\varepsilon}_t(\theta)) + \log p(\theta), \quad (4)$$

where  $\ell_t^K$  is the ROM1/Kalman contribution and  $\ell_{t,b}^{\text{obs}}$  measures observation fit at the recovered shock. The unconstrained code also adds the parameter-transform Jacobian. Because the shock is optimized rather than integrated out, this is a profiled observation-fit objective, not a marginal likelihood. I therefore report target means and target intervals rather than posterior moments.

## 4.4 The Smets–Wouters–HLT Model

I use *effective lower bound* (ELB) for the economic constraint and *occasionally binding constraint* (OBC) for the associated solution problem. I retain “ZLB” only in legacy artifact names or when describing a source that imposes a literal zero floor.

The model is Smets–Wouters (2007) in the Harding–Lindé–Trabandt (HLT) form. It places the quadratic investment adjustment cost from Section 3 inside a standard medium-scale model with habit, variable capital utilization, Kimball demand, sticky prices and wages, and a lower bound. The investment channel can therefore be compared with familiar nominal and real alternatives without changing models.

The HLT form sets the moving-average terms of the price- and wage-markup shocks to zero. I estimate seven shock persistences, seven shock volatilities, and four price and wage parameters, for 18 parameters in total. The online technical appendix gives the full model, priors, and computational record.

# 5 How Switching Fidelity Works

## 5.1 Residual Correction

ROM1 supplies a cheap first-order solution. At the same parameters, states, and shocks, SEP measures what that solution misses. A neural network learns the difference and adds it back only in stressed states. Formally,

$$\Delta_{\theta}(s, \varepsilon) = g_{\theta}^{\text{SEP}}(s, \varepsilon) - g_{\theta}^{\text{ROM1}}(s, \varepsilon), \quad \hat{g}_{\theta, w}(s, \varepsilon) = g_{\theta}^{\text{ROM1}}(s, \varepsilon) + \hat{\Delta}_w(s, \varepsilon, \theta).$$

The network targets  $\Delta_{\theta}$ , not  $g_{\theta}^{\text{SEP}}$ . This keeps the target small and makes validation direct: on held-out SEP solves, the error is exactly the error in the nonlinear correction.

I train the residual network offline on matched SEP and ROM1 solutions. During estimation, the inversion step recovers shocks, the gate turns on the residual in stressed periods,

and NUTS-HMC samples the parameters.

A preliminary chain also checks whether the training region is wide enough. If it reaches a stressed point outside that region, I solve the point with SEP and retrain before the final run. This avoids using the network as an extrapolator.

## 5.2 Stochastic Extended Path (SEP) Algorithm

SEP solves (1) directly after approximating expectations by Gauss–Hermite quadrature (Fair and Taylor, 1983; Adjemian and Juillard, 2025). The unknowns are all future states on a finite shock tree. Stacking the equations gives  $R(Y; \theta) = 0$ , solved by damped Newton steps,

$$Y^{(j+1)} = Y^{(j)} - \alpha^{(j)} \left[ \frac{\partial R}{\partial Y}(Y^{(j)}; \theta) \right]^{-1} R(Y^{(j)}; \theta), \quad (5)$$

where  $\alpha^{(j)}$  is chosen by line search. A single HLT SEP solve takes about 0.8 seconds, which is too slow to put inside every HMC leapfrog step but fast enough to generate an offline training set.

The lower bound is imposed inside the nonlinear system. For differentiability I use a smooth softplus version,

$$R_t = 1 + \frac{1}{\kappa_s} \log(1 + e^{\kappa_s(\tilde{R}_t - 1)}), \quad (6)$$

where  $\tilde{R}_t$  is the shadow Taylor-rule rate and  $\kappa_s = 100$ . Details on quadrature, warm starts, and solver tolerances are in the online technical appendix.

## 5.3 Dataset Generation and Surrogate Training

Training is supervised learning. I choose Sobol parameter draws, simulate state and shock inputs, solve SEP and ROM1 at the same inputs, and store

$$\Delta y_j = Hg_{\theta}^{\text{SEP}}(s_{j-1}, \varepsilon_j) - Hg_{\theta}^{\text{ROM1}}(s_{j-1}, \varepsilon_j).$$

The 18-parameter dataset contains 22,080 usable samples. Held-out non-binding RRMSE—the surrogate’s relative prediction error against the nonlinear SEP benchmark—is about 0.0012. The lower-bound stress region is less accurate and is treated separately. Architecture and training details are in the online technical appendix.

For each evaluated  $(s_{t-1}, \varepsilon_t, \theta)$ , I record feature  $z$ -scores. Gate-active points outside the checked range are sent back to SEP. The final training set combines the original design with these new points.

## 5.4 Local Solution Selection and Its Limit

The labels do more than train the network. They state which numerical solution the estimator is meant to follow. At every accepted input, the data contain one direct SEP transition and its residual relative to ROM1. A failed SEP solve is recorded as a failure; it is not replaced by a neural prediction. Within this audited domain, the procedure therefore enforces a unique supervised target.

I also perturb the target locally. The main ablation changes investment, utilization, price, and wage curvatures around 40 parameter vectors and repeats the SEP calculation at shock scales 0.1, 0.25, and 0.5. All paths converge. The support exercise separately records finite paths and SEP residuals for every promoted cell. Together these checks show that the selected branch is continuous over the states, parameters, and shock sizes used in the reported comparison. This is a numerical certificate of local branch consistency, not a theorem that the economic equilibrium is locally unique.

The distinction is important for nonlinear models. Local stability can hold inside a corridor even when larger shocks lead to another basin, a cycle, or an unstable path (Semmler and Sieveking, 1993; Schleer and Semmler, 2015). The present method does not enumerate roots, locate basin boundaries, trace bifurcations, or prove global stability. Hamiltonian Monte Carlo explores uncertainty conditional on the selected transition and initialized target basin; it cannot supply those missing equilibrium results. Global uniqueness and nonlinear

stability remain future work.

## 5.5 Regime-Switching Gate and Inversion Filter

Not every period needs the nonlinear correction. I classify a period as stressed when the state is far from steady state. Calm periods use ROM1; stressed periods use the residual-corrected transition,

$$s_t = (1 - m_t) g_{\theta}^{\text{ROM1}}(s_{t-1}, \varepsilon_t) + m_t \hat{g}_{\theta}(s_{t-1}, \varepsilon_t).$$

I recover shocks in two passes. The first pass solves a small regularized least-squares problem at each date. It is closed form under ROM1 and takes a few Gauss–Newton steps when the residual is active. The second pass holds those shocks fixed and samples only  $\theta$ . This avoids a particle filter and keeps the sampler at 18 dimensions.

The saving has a clear cost. Fixing shocks at their fitted values drops the shock density and the terms needed to integrate over shock uncertainty. I therefore compare profiled observation fit, not Bayes factors.

## 5.6 Filter-Free Hamiltonian Monte Carlo

I sample with NUTS-HMC ([Hoffman and Gelman, 2014](#)). In small validation models, shocks and the initial state can be sampled jointly with  $\theta$ . In the reported 18-parameter profiled runs, the inversion filter recovers shocks and NUTS samples  $\theta$  alone. A separate check samples shocks only in gate-active periods while holding calm periods on the ROM1 path. Four matched chains per branch permit rank, trace, effective-sample-size, Monte Carlo error, energy, tree-depth, and correlation diagnostics. Details are in the online technical appendix.

With the method in place, I turn to the economic question: where does first order fail, and what in the model makes it fail there?

## 6 Where the Nonlinearity Comes From

Throughout the comparison, the state, shock, and parameter vector are fixed; only the solution method changes. I call the resulting difference the SEP–ROM1 gap. It measures what first order misses within the region checked by the nonlinear solver.

### 6.1 How I Measure the Gap

For output  $k$  in sample  $j$ , write the difference as

$$d_{kj} = g_{k,\theta_j}^{\text{SEP}}(s_{j-1}, \varepsilon_j) - g_{k,\theta_j}^{\text{ROM1}}(s_{j-1}, \varepsilon_j).$$

To compare groups of variables, scale output  $k$  by a positive number  $a_k$  and compute

$$s_B(a) = \frac{\sum_j \sum_{k \in B} (d_{kj}/a_k)^2}{\sum_j \sum_k (d_{kj}/a_k)^2}.$$

This measure shows where the gap appears. The parameter changes below provide a separate check on the economic mechanism.

Applied to the seven series that enter estimation, the contrast is sharp. In model units, investment accounts for 78.4 percent of the squared SEP–ROM1 gap. After each series is scaled by its variation under SEP, the share is 66.2 percent. Inflation accounts for 0.17 and 0.34 percent under the same two measures. For the variables the model is asked to explain, first order therefore misses investment dynamics far more than price dynamics.

The gap has a natural economic interpretation. Capital utilization, investment adjustment costs, and Tobin’s  $q$  meet in the investment block. When a negative investment shock lowers the value of installed capital, the nonlinear solution retains the interactions among marginal adjustment costs, capital values, and discounting. ROM1 replaces them with fixed local coefficients.

The comparison matches inputs across solution methods and reports how the result

changes with scale. ROM1 is re-solved at each of the 120 parameter vectors in the source data. Stable first-order solutions exist for 117 vectors, leaving 21,528 observations in which both methods use the same state, shock, and parameters. The other three vectors are excluded rather than replaced or imputed.<sup>1</sup>

Table 1: Sensitivity of the parameter-matched SEP–ROM1 gap accounting

| Metric                             | Inv. (%) | Price (%) | Inv. rank | Largest group |
|------------------------------------|----------|-----------|-----------|---------------|
| Raw squared                        | 26.09    | 0.06      | 2         | Wage Phillips |
| Raw squared, equal variable weight | 26.29    | 0.05      | 2         | Wage Phillips |
| SEP-standardized squared           | 5.79     | 0.06      | 2         | Flex-Price    |
| Pooled-standardized squared        | 16.54    | 10.18     | 2         | Flex-Price    |
| Observables, raw squared           | 78.40    | 0.17      | 1         | Investment    |
| Observables, SEP-standardized      | 66.20    | 0.34      | 1         | Investment    |

*Notes:* SEP and ROM1 are evaluated at the same stored state, shock, and 18-parameter vector. The source contains 22,080 observations from 120 parameter vectors. Parameter-matched ROM1 solutions exist for 21,528 observations; 3 parameter vectors lack a stable first-order solution and are excluded from every row. Raw shares depend on the maintained model coordinates. SEP-standardized shares divide each variable’s gap by its sample standard deviation under SEP; pooled-standardized shares use the root mean of the SEP and ROM1 variances. Equal-variable weighting divides each block loss by its number of variables. Observable-only rows compare the seven measured series. These are descriptive locations of the transition discrepancy, not causal equation shares.

Scale matters because model variables have different units. Among the seven observables, model units and normalization by each series’ own variation both put investment first. A pooled normalization gives the series nearly equal weight and places investment fourth. Expanding the comparison to all 51 outputs brings in unobserved pricing, wage, and flex-price variables; investment and capital then rank second, with shares from 5.8 to 26.3 percent. The decomposition is therefore a statement about where the observable gap lies under the two reported series-level scales, not an invariant structural share.

The accounting locates the gap; changing the economic curvatures tests its source. I change the curvatures that govern investment, utilization, prices, and wages. I evaluate 40 posterior-region parameter vectors at shock scales 0.1, 0.25, and 0.5 and move each curvature by plus or minus 10 percent. Investment adjustment or capital utilization produces the

<sup>1</sup>The audit also reconstructs an earlier 68.9-percent calculation and finds that its stored first-order path held parameters at baseline while the nonlinear path varied them. I do not use that calculation. The replication record preserves the reconstruction and the reason for excluding it.

largest response for 38 of 40 parameter vectors at the smallest shock and for all 40 vectors at the two larger shocks. All paths converge. This test does not depend on how variables are assigned to blocks. It again locates the important nonlinearity on the real side of the model.

Figure 1 varies real-side and Kimball parameters around the baseline. Over the  $\pm 1$  standard-deviation grid, the real-side error range is 11.4 and the Kimball range is 1.4. Over the  $\pm 2$  grid, they are 23.5 and 2.6. Near the baseline, investment adjustment and utilization move the SEP-ROM1 error much more than Kimball curvature.

Pricing supplies a useful boundary case. At the high-Kimball calibration of [Harding et al. \(2022\)](#), one fixed-calibration raw accounting assigns 76.5 percent to investment and 2.0 percent to prices. More importantly, the local Kimball surface becomes steeper than its real-side counterpart. At  $\pm 2$  standard deviations, their ranges are 210.6 and 53.9. Thus pricing can dominate, but only after moving far from the curvature supported by the baseline and sampled parameter region.

## 6.2 Is the Lower Bound the Main Source?

A separate stress check asks whether the raw investment share falls as the lower bound starts to bind. Averaged over all simulated periods, investment accounts for 87.7–92.1 percent of the raw squared observable gap at shock scales from 0.8 to 1.5 (Figure 2). This is evidence about that raw metric on the plotted grid, not a general attribution of lower-bound effects.

Conditional on the rare binding periods, investment’s share does fall: it is 90 percent at scale 1.2 and 37 percent at scale 1.5. But the lower bound binds in only 0.1–0.4 percent of periods at the estimated parameters. The unconditional share therefore stays near 90 percent. The lower bound matters when it binds, but it does not drive the average gap in this experiment.

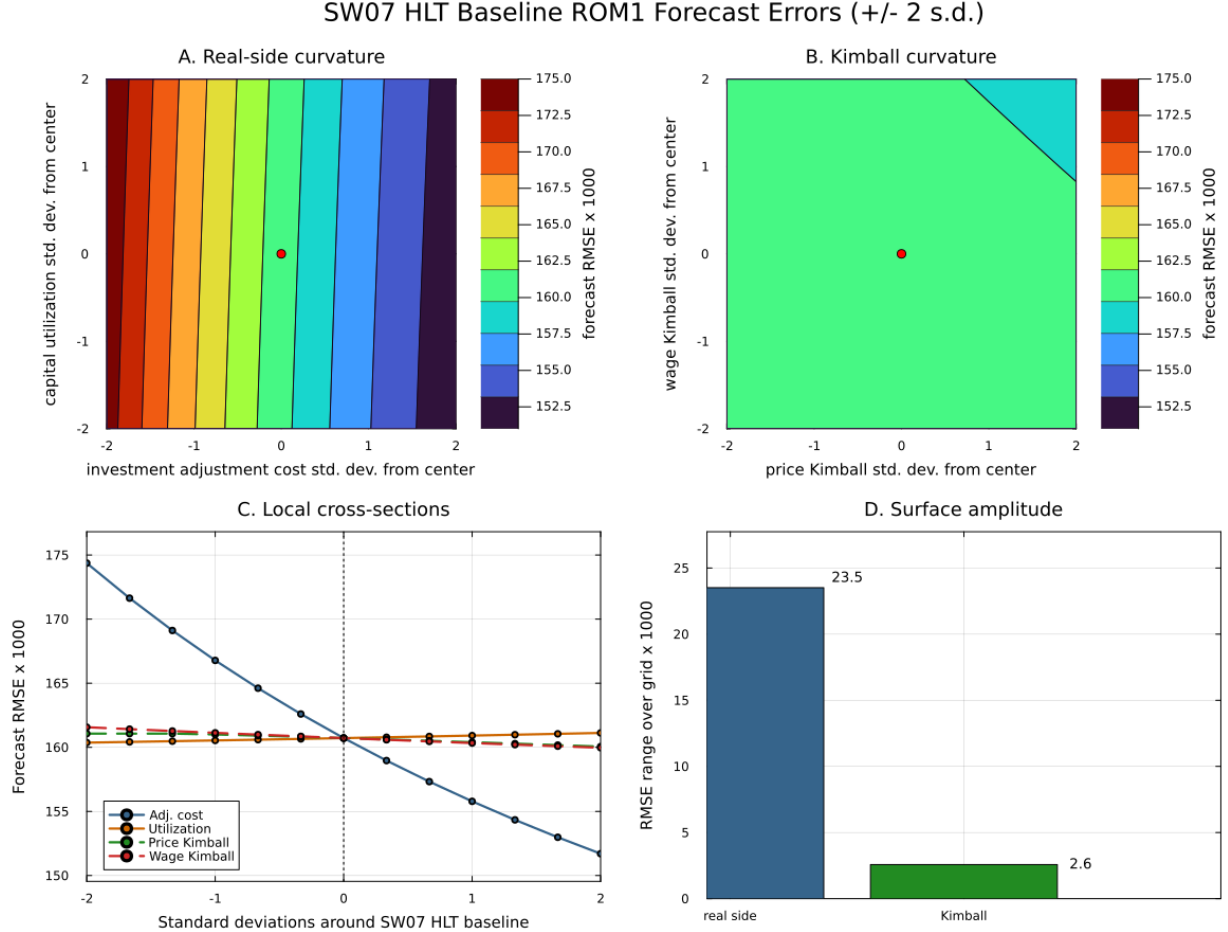


Figure 1: Baseline ROM1 forecast-error contours. The metric is 1,000 times the RMSE of one-step observable forecast errors, SEP minus ROM1, evaluated at the same SEP state, same shocks, and same parameter vector. Panels A and B plot filled contours over a  $\pm 2$  local-standard-deviation neighborhood around the maintained SW07–HLT baseline calibration; black contour lines are equal forecast-error levels, and both panels use the same color scale. The red marker denotes the center calibration. Panel C plots cross-sections through the center cell. Panel D reports the RMSE range over each grid. The real-side surface combines investment-adjustment and capital-utilization parameters; the Kimball surface combines price and wage curvature parameters.

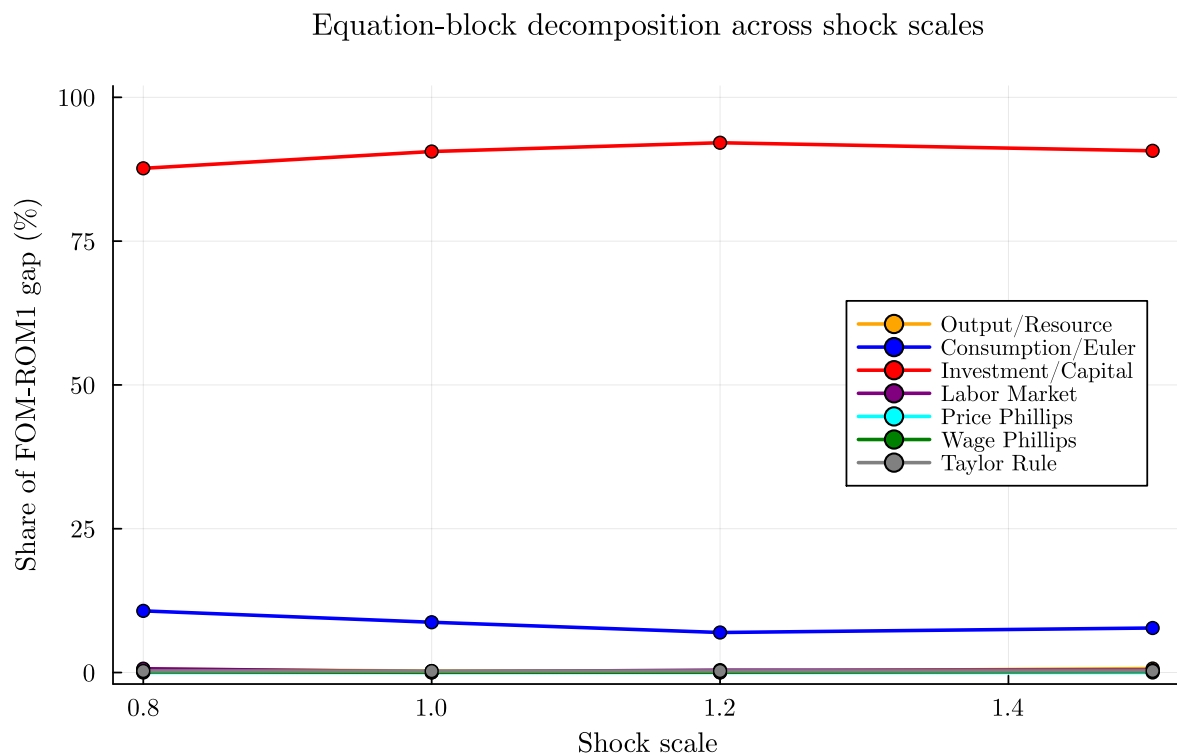


Figure 2: Observable shares as shock size rises from 0.8 to 1.5. Each point averages 15 parameter draws and 50 simulated periods per draw. Investment accounts for 87.7–92.1 percent of the raw squared SEP–ROM1 observable gap. Consumption accounts for 7.0–10.7 percent; every other observable remains below 1 percent. These raw-unit shares are not invariant to rescaling.

### 6.3 Can Higher-Order Local Solutions Close the Gap?

One might expect a higher-order perturbation to recover the investment nonlinearity without a direct pathwise solution. I test that possibility inside the same HLT model. I hold the published calibration, initial state, investment-specific shock, and twelve-quarter response window fixed. I then compare first-, pruned second-, and pruned third-order responses with a stochastic extended-path solution. I exclude the lower-bound equation from all four methods, so the comparison concerns the model’s smooth nonlinear dynamics.

Pruned second order helps, but only partly. It closes 4 percent of the scaled path gap at a half-standard-normal shock, 8 percent at a one-standard-normal shock, and 16 percent at a two-standard-normal shock. Pruned third order is farther from the direct path overall at all three sizes. Higher order can improve individual responses, including investment at the larger shocks, without bringing the joint path closer to the nonlinear solution.

Table 2: Higher-order local solutions against the nonlinear HLT path

| Method      | Shock size | Scaled RMSE | Investment RMSE | Gap closed |
|-------------|------------|-------------|-----------------|------------|
| ROM1        | 0.5        | 0.4234      | 0.0693          | 0.0%       |
| Pruned ROM2 | 0.5        | 0.4066      | 0.0721          | 4.0%       |
| Pruned ROM3 | 0.5        | 0.5554      | 0.0672          | -31.2%     |
| ROM1        | 1.0        | 0.4445      | 0.1370          | 0.0%       |
| Pruned ROM2 | 1.0        | 0.4084      | 0.1411          | 8.1%       |
| Pruned ROM3 | 1.0        | 0.5628      | 0.1296          | -26.6%     |
| ROM1        | 2.0        | 0.4887      | 0.3023          | 0.0%       |
| Pruned ROM2 | 2.0        | 0.4087      | 0.2752          | 16.4%      |
| Pruned ROM3 | 2.0        | 0.5783      | 0.2404          | -18.3%     |

*Notes:* All methods use the published HLT calibration, the same steady-state initial condition, the same investment-specific innovation, and twelve response quarters. The direct reference is stochastic extended path with a 24-quarter horizon. Scaled RMSE divides each response by its root mean square under the nonlinear reference before pooling variables. “Gap closed” is the reduction in this RMSE relative to ROM1 at the same shock size. The common benchmark uses the smooth HLT system without an active lower-bound equation.

This result changes the practical question. The relevant alternative to first order is not simply “use third order.” The approximation must be checked against the nonlinear path in the part of the model where the discrepancy appears. Switching fidelity uses that measured

discrepancy directly.

## 6.4 An Investment-Sensitive Policy Experiment

The solution gap matters for a policy comparison when policy changes the investment path. I feed the model an adverse one-standard-normal investment-specific shock and raise the Taylor rule's response to the output gap from 0.0593 to 0.25. All other parameters, the initial state, and the shock are unchanged. This is a transparent rule perturbation, not an optimal-policy or welfare calculation.

Under the nonlinear solution, the stronger response reduces the cumulative sixteen-quarter investment loss by 6.54 percentage-point quarters. ROM1 reports 5.39, about 18 percent less. The output effect goes in the opposite direction: ROM1 reports a slightly larger cumulative stabilization than the nonlinear solution, 2.64 against 2.35 percentage-point quarters. Thus first order gets the direction right but changes the measured strength and composition of the policy effect precisely where the earlier decomposition says it should.

## 6.5 Removing Investment Adjustment Costs

I also change the model itself. Full removal of investment adjustment costs makes the gap explode and sometimes causes SEP paths to fail. Because that intervention changes both equilibrium dynamics and solver support, I treat it as a stress test, not as an additive mechanism decomposition. The posterior-region plus-or-minus-10-percent exercise above is the maintained local evidence.

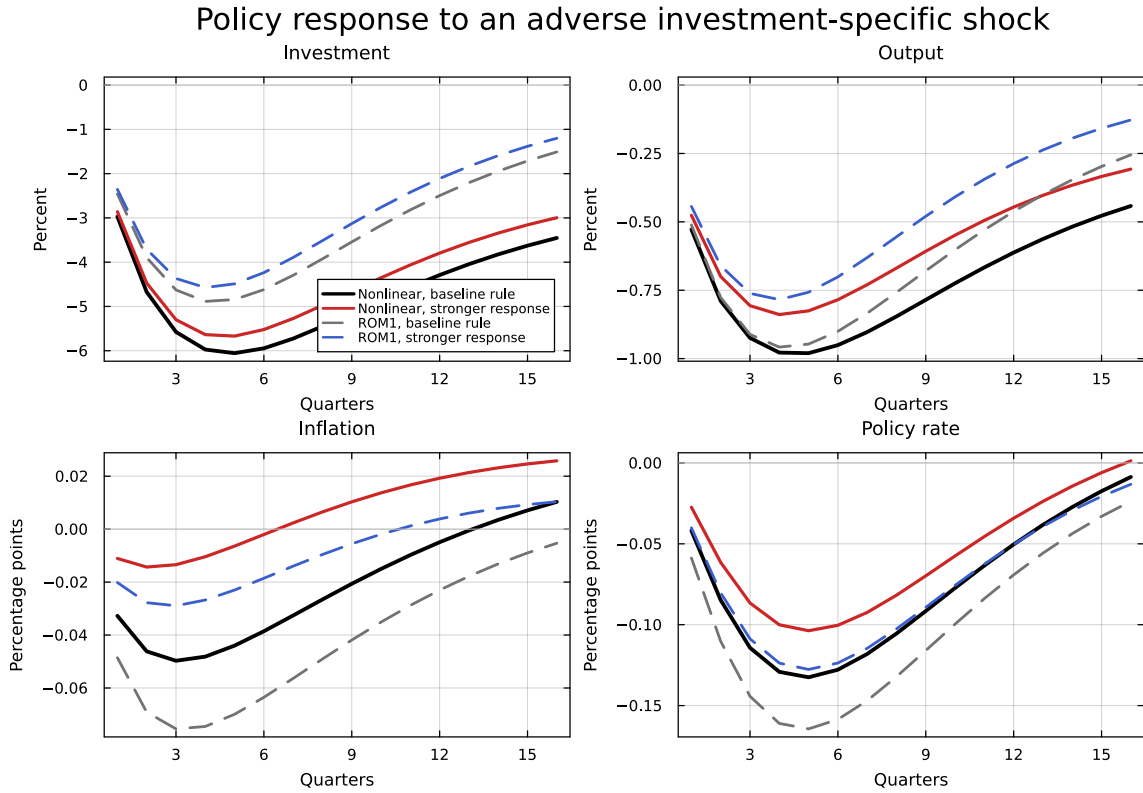


Figure 3: Response to an adverse investment-specific shock under the baseline Taylor rule and a stronger output response. Solid lines are stochastic extended-path solutions; dashed lines are ROM1. Each pair holds the model, initial state, and shock fixed. The exercise reports transition paths, not welfare or an optimal rule.

## 7 What Changes in Estimation

There are two separate questions. First, does the neural correction recover what ungated ROM1 misses? This requires direct SEP, ROM1, and ROM1 plus the correction at the same inputs. Second, what can be learned about parameters once the correction enters a profiled 18-parameter target? This is where HMC is useful. Keeping the questions separate avoids treating a gated estimator as if it were plain ROM1.

### 7.1 The Correction Relative to Parameter-Matched ROM1

The clean comparison uses direct SEP paths from 120 draws of all 18 estimated parameters. ROM1 is solved again at every draw. Three parameter vectors do not have a stable first-order solution and are recorded as failures. From the 117 common-support vectors, I reserve twelve complete vectors for validation and train on the other 105. The validation set therefore contains 2,208 transitions at parameter values the network never sees during training.

At each transition, direct SEP, ungated ROM1, and ROM1 plus the learned residual use the same state, shock, and parameter vector. The gate and inversion filter do not enter. Raw RMSE pooled across the seven observables falls from 2.188 under ROM1 to 1.237 after the correction, a 43.4 percent reduction. Every observable improves, although the gains range from 2.1 percent for the policy rate to 60.6 percent for wage growth.

This is the relevant test of the correction. It is harder than a random transition split because no validation parameter vector appears in training, but it remains a finite normal-state corridor rather than a direct nonlinear posterior. The earlier 10,000-cell four-parameter bridge, reported in the online appendix, uses ROM1 fixed at its baseline calibration. It is useful as a composite-map emulation check, not as evidence for a parameter-matched residual.

Table 3: Parameter-held-out correction relative to ungated ROM1

| Observable         | ROM1 RMSE | Corrected RMSE | Reduction |
|--------------------|-----------|----------------|-----------|
| Output growth      | 0.7727    | 0.3950         | 48.9%     |
| Consumption growth | 2.3068    | 1.1050         | 52.1%     |
| Investment growth  | 5.2089    | 3.0128         | 42.2%     |
| Hours              | 0.3530    | 0.3226         | 8.6%      |
| Inflation          | 0.2829    | 0.2127         | 24.8%     |
| Wage growth        | 0.3973    | 0.1565         | 60.6%     |
| Policy rate        | 0.3051    | 0.2987         | 2.1%      |
| Raw pooled RMSE    | 2.1879    | 1.2373         | 43.4%     |

*Notes:* ROM1 is re-solved at each 18-parameter vector. The split holds out twelve complete parameter vectors, or 2,208 transitions; 105 vectors, or 19,320 transitions, are used for training. Direct SEP, ROM1, and the corrected transition share the same states, shocks, and parameters. RMSE is reported in the maintained observable coordinates, so the pooled row is not invariant to rescaling.

## 7.2 What HMC Reveals

The 18-parameter exercise has four chains, each with 500 warmup and 1,000 retained draws. HMC makes the joint target visible in ways that point estimates cannot. It provides marginal intervals, cross-parameter correlations, effective sample sizes, Monte Carlo standard errors, rank and trace checks, and diagnostics of energy exploration and trajectory length.

Table 4: Four-chain HMC diagnostics for the profiled targets

| Diagnostic                   | ROM1 plus NN | Linear plus gate |
|------------------------------|--------------|------------------|
| Maximum $\hat{R}$            | 1.0093       | 1.0026           |
| Minimum bulk ESS             | 653          | 1741             |
| Minimum tail ESS             | 976          | 1211             |
| Maximum MCSE / marginal SD   | 0.037        | 0.023            |
| Minimum chain E-BFMI         | 0.784        | 0.851            |
| Divergences                  | 2            | 0                |
| Retained numerical errors    | 2            | 0                |
| Mean depth-cap share         | 79.1%        | 0.0%             |
| Mean acceptance              | 0.888        | 0.920            |
| Largest absolute correlation | 0.541        | 0.502            |

*Notes:* Each branch contains 4,000 retained draws. E-BFMI is computed chain by chain. The depth-cap share is the fraction of draws that reach the maintained maximum tree depth of 8. The diagnostics assess the profiled target within the initialized basin; they neither normalize the target as a marginal likelihood nor search for other economic solution branches.

For the corrected target, maximum  $\hat{R}$  is 1.0093, minimum bulk ESS is 653, minimum tail ESS is 976, and the largest Monte Carlo standard error is 3.7 percent of its marginal standard deviation. Chain E-BFMI ranges from 0.784 to 0.877. The strongest posterior correlation is  $-0.54$ , between price-markup persistence and its innovation scale. These results support within-basin estimation of the reported marginals.

They also reveal a difficult target. Two of 4,000 corrected draws diverge, and 79.1 percent reach the depth cap. The problem is therefore trajectory length, not a failure of the chains to overlap. The full parameter table, trace plots, rank plots, autocorrelation functions, interval plot, energy diagnostics, and correlation matrix are reported in the online technical appendix. None of these objects proves that the economic model has a globally unique or stable nonlinear solution.

### 7.3 An Implementation Ablation Inside the HMC Target

I also run the production sampler with and without its maintained neural residual. The comparison branch keeps the regime switch and inversion filter, so it is linear plus gate rather than plain ROM1. Across four matched seeds, the residual changes profiled observation fit by only  $-0.102$  nat on average. That result is an implementation check, not the correction benchmark above. The stored production residual was trained relative to baseline ROM1 and then added to ROM1 re-solved at each proposal; it is therefore not the parameter-matched training design in Table 3. The online appendix reports the eight chains and the full comparison.

### 7.4 Computation

The nonlinear labels cost 10 hours to generate, and training takes 2 hours on an Apple M4. These are one-time costs. The residual does not depend on the observed sample, so the same trained network can be used in later estimation runs as long as their inputs remain inside the checked region. Using the residual inside HMC is about 2,450 times faster than solving

SEP at every target evaluation.

## 8 Robustness and Scope

I use the robustness checks to ask when the investment result holds and when it can fail. I change the pricing technology, vary investment-adjustment curvature, and isolate periods when the lower bound binds. These exercises keep the raw observable measure fixed so that only the economic specification changes.

### 8.1 Does the Result Depend on Price Setting?

In raw units, the matched baseline audit assigns 0.06 percent of the full-output gap to the price block and 0.17 percent of the observable gap to inflation. I check whether that raw-unit result is specific to Calvo pricing. The alternatives are a high-slope Calvo case, a moderate-slope Calvo case, and Rotemberg pricing. The investment, wage, and policy blocks remain unchanged.

Investment accounts for 89–91 percent of the gap in the two Calvo variants. The price share rounds to zero. Rotemberg pricing is the stronger test because its quadratic price-adjustment cost is visible at second order. Even there, investment accounts for 83–86 percent and prices for less than 0.5 percent. Inflation remains close enough to trend that the visible price curvature has a small quantitative effect.

### 8.2 Does Adjustment-Cost Curvature Change the Result?

The baseline sets  $S''(1) = 6$ . I repeat the decomposition at values from 2 to 10, using 25 parameter draws and 40-period SEP paths at each value.

Investment and capital account for 89–95 percent of the gap across this range. The price Phillips curve rounds to zero at every value. Higher curvature reduces the absolute gap from 5.6 to 2.8 model units because it restrains investment more strongly. Thus curvature changes

the size of the gap more than its composition in this design.

Table 5: Raw observable-gap sensitivity to  $S''(1)$

| $S''(1)$   | Investment | Consumption | Output | Labor | Price | Wage | Taylor | Gap    |
|------------|------------|-------------|--------|-------|-------|------|--------|--------|
| 2.0        | 95.1       | 4.5         | 0.2    | 0.1   | 0.0   | 0.0  | 0.1    | 5.6044 |
| 4.0        | 94.1       | 5.5         | 0.1    | 0.1   | 0.0   | 0.0  | 0.2    | 3.6844 |
| <b>6.0</b> | 92.2       | 7.3         | 0.1    | 0.1   | 0.0   | 0.0  | 0.3    | 3.6558 |
| 8.0        | 91.8       | 7.5         | 0.1    | 0.2   | 0.0   | 0.0  | 0.3    | 2.7908 |
| 10.0       | 88.8       | 10.2        | 0.3    | 0.3   | 0.0   | 0.0  | 0.4    | 2.7618 |

*Notes:* Each column reports its share of the raw squared FOM–ROM1 gap across the seven observables. “Gap” is the root-mean-square observable gap in model units. The baseline calibration  $S''(1) = 6.0$  is shown in bold. Shares sum to 100% by construction, but they depend on observable units and are not orthogonal or causal equation contributions. The exercise uses 25 representative parameter vectors and 40 periods per trajectory at unit shock scale.

### 8.3 Scope of the Evidence

The result is local in both parameters and states. The matched decomposition covers 117 parameter vectors with stable first-order solutions, and the residual validation holds out twelve of them. The separate ablation covers 40 nearby vectors at three shock sizes. These designs make the tested region visible; they do not turn it into global support.

Supervised labels provide one direct numerical target at each accepted input. That is useful because the network cannot quietly choose a different branch or replace a failed solve with its own equilibrium loss. It is still a local certificate. The exercises do not establish a unique target basin, a full direct-SEP posterior, or accuracy after the economy leaves the labeled region.

Global stability is a further limitation. The model may have another solution outside the audited corridor, and a large enough shock may move the economy to another basin or an unstable path. The present SEP-labeling and HMC procedure does not map those possibilities. Addressing them requires a separate global solution and stability analysis.

The transition-gap shares are likewise tied to units, scale choices, and the set of variables. The matched audit reports several normalizations rather than declaring one structurally

correct. It also excludes three of 120 parameter vectors for which ROM1 has no stable first-order solution. Those failures mark a limit of the linear approximation, but the audit cannot say how another solution concept would change the shares at those points.

## 8.4 What the Result Means

The result does not calculate welfare, but it does show where a first-order policy analysis is most likely to miss an important response: investment, capital, and their effects on output and hours. When a policy experiment moves investment or the value of installed capital, the real-side dynamics deserve a nonlinear check before price nonlinearities are treated as the main concern.

## 9 Conclusion

First order does not fail evenly in the Smets–Wouters model. Around the estimated region it tracks price dynamics closely but misses much of the investment response. The reason lies in the adjustment cost and in its interaction with capital values and discounting, terms that disappear in a first-order expansion. The same real-side mechanism appears in the observable gaps and in the curvature experiments. Price curvature takes over only in deliberately extreme cases.

Higher-order local solutions do not remove the problem. In the common investment-shock experiment, pruned second order closes at most 16 percent of the scaled path gap, while pruned third order is farther from the direct path overall. The policy experiment shows why that remaining difference can matter: first order understates the investment stabilization from a stronger output response by about 18 percent.

Switching fidelity makes the direct comparison affordable without severing it from the nonlinear model. Direct nonlinear solutions measure the error of first order offline, and a neural network learns that measured error. When ROM1 is re-solved at each parameter vector and twelve complete vectors are held out, the correction lowers raw pooled observable RMSE by 43.4 percent and improves all seven observables. The gate and inversion filter play no role in this accuracy test.

Supervised SEP labels make the local solution choice visible. The parameter ablations and three shock scales stay on the same convergent branch, and failed support points are recorded rather than imputed. The method nevertheless follows one local branch. It neither proves global uniqueness nor maps stability outside the audited corridor.

The lesson is not that nominal rigidities are unimportant. It is that, when this model moves away from steady state, the first place to look for nonlinear dynamics is investment. Policy exercises that move investment, capital utilization, or the value of installed capital should check that margin with a nonlinear solution. Switching fidelity provides a practical and testable way to do so within a clearly stated local domain.

## References

- Stéphane Adjemian and Michel Juillard. Stochastic extended path. Working paper, Dynare Team, 2025. URL <https://stephane-adjemian.fr/papers/sep-2025.pdf>. Accessed: 2026-02-27.
- Martin M. Andreasen, Jesús Fernández-Villaverde, and Juan F. Rubio-Ramírez. The pruned state-space system for non-linear dsge models: Theory and empirical applications. *Review of Economic Studies*, 85(1):1–49, 2018. doi: 10.1093/restud/rdx037.
- Christophe Andrieu, Arnaud Doucet, and Roman Holenstein. Particle markov chain monte carlo methods. *Journal of the Royal Statistical Society: Series B*, 72(3):269–342, 2010. doi: 10.1111/j.1467-9868.2009.00736.x.
- S. Borağan Aruoba, Pablo Cuba-Borda, Kenji Higa-Flores, Frank Schorfheide, and Sergio Villalvazo. Piecewise-linear approximations and filtering for DSGE models with occasionally binding constraints. *Review of Economic Dynamics*, 41:96–120, 2021. doi: 10.1016/j.red.2020.12.003.
- Marlon Azinovic, Luca Gaegauf, and Simon Scheidegger. Deep equilibrium nets. *International Economic Review*, 63(4):1471–1525, 2022. doi: 10.1111/iere.12575.
- David Childers, Jesús Fernández-Villaverde, Jesse Perla, Christopher Rackauckas, and Peifan Wu. Differentiable state-space models and hamiltonian monte carlo estimation. NBER Working Paper 30573, National Bureau of Economic Research, 2022.
- Nicolas Chopin, Pierre E. Jacob, and Omiros Papaspiliopoulos. SMC<sup>2</sup>: An efficient algorithm for sequential analysis of state space models. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 75(3):397–426, 2013. doi: 10.1111/j.1467-9868.2012.01046.x.
- Ray C. Fair and John B. Taylor. Solution and maximum likelihood estimation of dynamic

nonlinear rational expectations models. *Econometrica*, 51(4):1169–1185, 1983. doi: 10.2307/1912057.

Mátyás Farkas and Bálint Tatár. Bayesian estimation of DSGE models with Hamiltonian Monte Carlo. IMFS Working Paper Series 144, Goethe University Frankfurt, Institute for Monetary and Financial Stability (IMFS), 2020. URL <https://www.econstor.eu/bitstream/10419/223402/1/1728957516.pdf>. Adapts HMC to DSGE estimation via Stan; documents a flat Smets–Wouters posterior target associated with markup ARMA dynamics.

Jesús Fernández-Villaverde and Oren Levintal. Solution methods for models with rare disasters. *Quantitative Economics*, 9(2):903–944, 2018. doi: 10.3982/QE744.

Jesús Fernández-Villaverde and Juan F. Rubio-Ramírez. Estimating macroeconomic models: A likelihood approach. *Review of Economic Studies*, 74(4):1059–1087, 2007. doi: 10.1111/j.1467-937X.2007.00437.x.

Luca Guerrieri and Matteo Iacoviello. OccBin: A toolkit for solving dynamic models with occasionally binding constraints easily. *Journal of Monetary Economics*, 70:22–38, 2015. doi: 10.1016/j.jmoneco.2014.08.005.

Martín Harding, Jesper Lindé, and Mathias Trabandt. Resolving the missing deflation puzzle. *Journal of Monetary Economics*, 126:15–34, 2022. doi: 10.1016/j.jmoneco.2021.09.003.

Edward P. Herbst and Frank Schorfheide. Tempered particle filtering. *Journal of Econometrics*, 210(1):26–44, 2019. doi: 10.1016/j.jeconom.2018.11.003.

Matthew D. Hoffman and Andrew Gelman. The no-u-turn sampler: Adaptively setting path lengths in hamiltonian monte carlo. *Journal of Machine Learning Research*, 15(1):1593–1623, 2014.

- Eva F. Janssens and Sean McCrary. Efficient estimation of nonlinear DSGE models. Working Paper 5282668, SSRN, 2025. URL [https://papers.ssrn.com/sol3/papers.cfm?abstract\\_id=5282668](https://papers.ssrn.com/sol3/papers.cfm?abstract_id=5282668).
- Kenneth L. Judd. Projection methods for solving aggregate growth models. *Journal of Economic Theory*, 58(2):410–452, 1992. doi: 10.1016/0022-0531(92)90061-L.
- Hanno Kase, Leonardo Melosi, and Matthias Rottner. Estimating nonlinear heterogeneous agent models with neural networks. BIS Working Paper 1241, Bank for International Settlements, 2025. URL <https://www.bis.org/publ/work1241.htm>.
- Jinill Kim, Sunghyun Kim, Ernst Schaumburg, and Christopher A. Sims. Calculating and using second-order accurate solutions of discrete time dynamic equilibrium models. *Journal of Economic Dynamics and Control*, 32(11):3397–3414, 2008. doi: 10.1016/j.jedc.2008.02.003.
- Robert Kollmann. Tractable likelihood-based estimation of non-linear DSGE models. *Economics Letters*, 161:90–92, 2017. doi: 10.1016/j.econlet.2017.08.027.
- Felix Kubler. Verifying competitive equilibria in dynamic economies. *Review of Economic Studies*, 78(4):1379–1399, 2011. doi: 10.1093/restud/rdr005.
- Lilia Maliar, Serguei Maliar, and Pablo Winant. Deep learning for solving dynamic economic models. *Journal of Monetary Economics*, 122:76–101, 2021. doi: 10.1016/j.jmoneco.2021.07.004.
- Benjamin A. Malin, Dirk Krueger, and Felix Kubler. Solving the multi-country real business cycle model using a smolyak-collocation method. *Journal of Economic Dynamics and Control*, 35(2):229–239, 2011. doi: 10.1016/j.jedc.2010.09.015.
- Alessandro Mennuni, Juan F. Rubio-Ramírez, and Serhiy Stepanchuk. Dynamic perturbation. *Review of Economic Studies*, 92(2):1157–1192, 2025. doi: 10.1093/restud/rdae037.

- Christopher Naubert. Differentiable, filter-free bayesian estimation of DSGE models using mixture density networks. Staff Working Paper 2025-3, Bank of Canada, January 2025. URL <https://www.bankofcanada.ca/2025/01/staff-working-paper-2025-3/>.
- Benjamin Peherstorfer, Karen Willcox, and Max Gunzburger. Survey of multifidelity methods in uncertainty propagation, inference, and optimization. *SIAM Review*, 60(3):550–591, 2018. doi: 10.1137/16M1082469.
- Frauke Schleer and Willi Semmler. Financial sector and output dynamics in the euro area: Non-linearities reconsidered. *Journal of Macroeconomics*, 46:235–263, 2015. doi: 10.1016/j.jmacro.2015.09.001.
- Stephanie Schmitt-Grohé and Martín Uribe. Solving dynamic general equilibrium models using a second-order approximation to the policy function. *Journal of Economic Dynamics and Control*, 28(4):755–775, 2004. doi: 10.1016/S0165-1889(03)00043-5.
- Willi Semmler and Malte Sieveking. Nonlinear liquidity-growth dynamics with corridor-stability. *Journal of Economic Behavior & Organization*, 22(2):189–208, 1993. doi: 10.1016/0167-2681(93)90016-S.
- Chris Snyder, Thomas Bengtsson, Peter Bickel, and Jeff Anderson. Obstacles to high-dimensional particle filtering. *Monthly Weather Review*, 136(12):4629–4640, 2008. doi: 10.1175/2008MWR2529.1.