

Slow Learning

Lawrence J. Christiano Martin Eichenbaum
Benjamin K. Johannsen*

May 21, 2024

Abstract

This paper investigates what features of an economy determine whether convergence under learning is fast or slow. In all of the models that we consider, people's beliefs about model outcomes are central determinants of those outcomes. We argue that under certain circumstances, convergence of a learning equilibrium to the rational expectations equilibrium can be so slow that policy analysis based on rational expectations is very misleading. We also develop new analytic results regarding rates of convergence in learning models.

*The analysis and conclusions set forth are those of the authors and do not indicate concurrence by the Board of Governors or anyone else. The authors thank Klaus Adam, George-Marios Angeletos, Ed Herbst, Albert Marcet, Michael Massman, Alexander Mayer, Damjan Pfajfar, Bruce Preston, Tom Sargent, Robert Tetlow, Ivan Werning, Mike Woodford and participants in a number of seminars for helpful comments and suggestions.

1 Introduction

Rational expectations may be a useful modeling strategy in tranquil times like the Great Moderation. This strategy is less appealing, however, when people are confronted with novel events, such as the Great Recession or the COVID-19 pandemic. This fact was self evident to the founders of rational expectations. For example, referring to the model in his seminal paper on asset prices, Robert E. Lucas, Jr. writes:

....The model described above “assumes” that agents know a great deal about the structure of the economy, and perform some non-routine computations. It is in order to ask, then: will an economy with agents armed with “sensible” rules-of-thumb, revising these rules from time to time so as to claim observed rents, tend as time passes to behave as described...”

Lucas (1978, p. 1437)

This paper analyzes the evolution of economic aggregates after a novel event. We assume that people must learn about their environment by forming beliefs about future economic outcomes and update those beliefs as the data come in.

There is a voluminous literature that addresses the question of whether economies in which people are learning about their environment converge to a rational expectations equilibrium (REE).¹ In contrast, much less attention has been paid to the question of how long it takes to converge to an REE.²

The answer to this question is critical to assessing the usefulness of rational expectations for understanding the effects of shocks on the economy and the efficacy of policies to deal with the repercussions of those shocks. As Vives (1993, p. 329) writes, in a changing world, for all practical purposes, “‘slow’ convergence may mean no convergence.” This paper investigates what features of an economy determine whether learning is fast or slow. Critically, in all of the models we consider, people’s beliefs about model outcomes are central determinants of equilibrium outcomes.³

Our central finding is that when beliefs are partially self-fulfilling, learning equilibria

¹See, for example, Bray and Savin (1986), Marcet and Sargent (1989b), and Evans and Honkapohja (2001).

²Some exceptions include Vives (1993), Marcet and Sargent (1995), Ferrero (2007), and Chien et al. (2021).

³There is a large literature that explores the speed with which people learn the parameters of *exogenous* stochastic processes. For example, Erceg and Levin (2003), Gust et al. (2018) and Farmer et al. (2021) describe an empirically relevant set of time series representations with hard-to-learn low frequency components.

converge slowly to rational expectations. Indeed, learning can be extraordinarily slow, with progress being measured in *millennia*. Under these circumstances, policy analyses based on rational expectations can be very misleading.

We begin by considering a model developed by Bray and Savin (1986) that is a workhorse in the learning literature. An important virtue of the model is that it is a very simple setup in which people’s beliefs determine market outcomes. The reduced form of the model encompasses cases in which beliefs are partially self-fulfilling and cases in which beliefs are self-defeating. A particular parameter, which we denote by b , controls how beliefs about market outcomes affect actual market outcomes. When $0 < b < 1$, beliefs are partially self fulfilling. The Lucas (1973) supply model, in which a higher expected price level leads to a higher actual price level, falls into this case. When $b < 0$, beliefs are self defeating. Muth (1961)’s version of the classic Cobweb model, in which a higher expected price level leads to a lower actual price level, falls into this case.

In the REE of the Bray and Savin (1986) model, people’s beliefs about economic aggregates are constants, independent of past data. Bray and Savin (1986) show that when people behave like Bayesians, their beliefs converge almost surely to rational expectations. In contrast, we focus on the *rate* at which beliefs converge under both Bayesian and classical (least squares) learning. In both cases, beliefs are a stochastic process that depends on past data. Building on Ljung (1977), Marcet and Sargent (1989c) and Evans and Honkapohja (2001) show that in a class of learning models, the eigenvalues of a particular ordinary differential equation (ODE) determine whether or not the system converges asymptotically. Those eigenvalues can also be used to characterize the amount of time it takes for learning to converge. In the Bray and Savin (1986) model, there is only one eigenvalue and it is equal to b . Convergence can be very slow, depending on the value of b .

Consistent with results in Christopheit and Massmann (2018), for moderately high values of b , we show that it takes an extraordinarily large number of periods to close two-thirds of the expected gap between the initial priors and the rational expectations belief. The intuition behind the possibility of slow convergence is as follows. When people’s expectations of a variable are largely self-fulfilling, they are slow to adjust their priors, and it takes a long time for them to converge to rational expectations. In contrast, when people’s expectations lead to outcomes different from their beliefs, they are quick to change those beliefs, and convergence is fast. For convenience, we refer to

this intuition as the *learning principle*.⁴

The possibility of slow convergence is not just a theoretical curiosum. The binding zero lower bound (ZLB) during the Great Recession was a novel event that few people understood when it first occurred. We argue that when the ZLB is binding, learning is particularly slow. We also argue that the implications of slow learning for policy are substantial: The predicted effects of monetary and fiscal policies are very different under rational expectations than under slow learning. We make these arguments using a simple version of the New Keynesian (NK) model that was widely used to understand the Great Recession and the effect of the binding ZLB.⁵

We begin by considering learning equilibria in the ZLB in the absence of government interventions. Our key result is that convergence is very slow. Indeed, in the benchmark parameterization of the model the ZLB would almost certainly be over long before learning has come close to converging. The reason for slow convergence is that in the ZLB, the expectations of households and firms tend to be self-fulfilling. To understand why, suppose that firms and households expect lower inflation in the future. Because of price-setting adjustment costs, firms are incentivized to cut prices today. In the ZLB, low inflation expectations mean households believe the real interest rate is high. Consequently, households reduce their demand for consumption, which leads to a fall in the marginal cost of production. Hence, the actions of both households and firms lead to lower current inflation. With learning, low current inflation shifts expected inflation down in the next period. The previous mechanism repeats itself in the next period so that actual inflation in the next period is also low. We conclude that, in the ZLB, deflation expectations are partially self fulfilling, and the NK model behaves like a high b economy.

In the NK model, when people have rational expectations, a shock that triggers a binding ZLB leads to a sharp decline in inflation and output (see Eggertsson and Woodford (2004)). The large effects arise because the shock triggers high expected deflation and real interest rates. In contrast, under learning the same shock leads only to a moderate and gradual decline in inflation because, with learning, expectations are partially backward-looking. So, if people begin the episode not expecting a large deflation, then the actual decline in inflation will be relatively moderate.

⁴A version of the learning principle has been studied by Heemeijer et al. (2009) and Hommes (2011).

⁵Much of the work in the initial aftermath of that event combined rational expectations with the NK model. See, for example, Eggertsson and Woodford (2004), Christiano et al. (2011) and Del Negro et al. (2023).

Next, we consider the effects of fiscal policy in the ZLB. We find that the efficacy of fiscal policy is much smaller under learning than under rational expectations. Under rational expectations, the multiplier is very large in the ZLB because an increase in government purchases causes a rise in expected inflation (see Christiano et al. (2011)). Because the nominal interest rate is fixed, this rise generates a fall in the real interest rate, a rise in consumption, and a multiplier substantially larger than unity. Under learning, expected inflation is partially backward-looking and does not move much after an increase in government purchases. So, the real interest does not fall by very much, and consumption rises by only a small amount. As a result, the key driver of the large REE multiplier is effectively eliminated, and the multiplier is close to unity.

Next, we turn to the efficacy of monetary policy in the wake of a shock to the discount rate, under learning. We begin by considering the effects of a simple form of forward guidance: The monetary authority commits to keeping the nominal interest rate at zero for one period after the shock that makes the ZLB binding returns to its steady-state level. Interestingly, the number of REEs proliferates under forward guidance. However, we show that only one equilibrium is stable under learning. Consistent with the existing literature (for example, Del Negro et al. (2023) and Woodford (2012)), we find that forward guidance is powerful under rational expectations. As is well-known, the power of forward guidance under rational expectations reflects its strong effect on expected inflation. Under learning, the effects of forward guidance have very little influence on expected inflation because expectations are partially backward-looking. It follows that under learning, forward guidance is not very powerful. So, as with fiscal policy, a rational expectations-based analysis of monetary policy can be very misleading.

In our analysis, people fully integrate the fact that they are learning when they solve their problems. For convenience, we refer to this approach as *internalized learning*. We formulate households' and firms' problems in recursive form. Given the recursive structure of learning, this approach seems natural: People start a period with an initial set of beliefs, then see data and update their beliefs using Bayes' rule. In our environment, aggregate prices adjust in a given period to clear markets. In our learning equilibrium, however, we do not require that planned future individual decisions are market clearing or that the sum of expected future individual decisions coincides with the corresponding expected aggregate outcomes. People's value functions incorporate their understanding that they will continue learning and adapting their behavior as

new data arrive.

As it turns out, implementing internalized learning in the nonlinear solution of the model is computationally very challenging. In a learning equilibrium, the parameters that characterize beliefs are also state variables and this greatly exacerbates the curse of dimensionality.⁶ See Appendix C for details.

Next, we study the asymptotic properties of the model by linearizing its solution. We establish three sets of results. First, we identify the analog of b in the linearized solution, which determines the asymptotic rate of convergence of the NK model. It is the largest real part of the eigenvalues of the matrix that maps beliefs about the state of the economy into their realized values. This property is not specific to the simple NK model. It applies to a broader class of multivariate models and extends some results in Christopheit and Massmann (2018) to those multivariate models. So our result about the analog of b allows an analyst working with a model falling within that class to quickly determine whether learning is slow or fast. Second, we show that the linearized solution to the NK is a good approximation to the nonlinear solution in a nontrivial neighborhood about the point at which the approximation is taken. Third, we show that the asymptotic rate of convergence in mean beliefs is a good guide to the small t (finite sample) rate of convergence.

In contrast to internalized learning, much of the learning literature works with a version of Kreps (1998)'s *Anticipated Utility* approach. In this approach, people update their beliefs every period as new data come in. But, when they make their decisions, people proceed as though their beliefs will never be revised again. This approach has been criticized for its internal inconsistency (see Cogley and Sargent (2008) and Adam and Marcet (2011)).

We simulate our model under both internalized learning and anticipated utility. We find that the results under both approaches are qualitatively similar. However, for some experiments there are important quantitative differences between the two approaches due to the more prominent role played by uncertainty under internalized learning. These results are consistent with those obtained by Cogley and Sargent (2008), who studied a stochastic endowment economy with a storage technology.

The remainder of this paper is organized as follows. Section 3 analyzes learning in the Bray and Savin (1986) environment. Section 4 discusses our approach to learning in the NK model. Section 5 characterizes the set of minimal state variable REE while

⁶We solve our model using a compiled programming language (c++) and we make use of more than 300 processors.

the representative household's discount rate is low. Section 6 analyzes the local and global learnability of those equilibria. In Section 7, we analyze the speed of convergence of the learning equilibrium in the NK model after a drop in the discount rate. In that section, we also compare the internalized learning and anticipated utility approaches to learning. Section 8 assesses the sensitivity of the efficacy of fiscal policy and forward guidance to learning. Section 9 extends our analytic results about rates of convergence reported in section 9 to the vector case that encompasses the NK model. Section 10 contains concluding remarks.

2 Related Literature

Our paper relates to a number of literatures. The first is the literature that studies the properties of recursive stochastic estimators in learning models. As noted above, Ljung (1977) establishes that a recursive estimator, $\hat{\theta}_t$, converges almost surely to a limiting value, θ , if a particular ordinary differential equation, ODE, which is determined by the underlying system, has eigenvalues with real parts that are less than unity. (1989c; 1989a), Woodford (1990), (2000; 2001) and others build on Ljung (1977) to study the conditions under which learning equilibria converge to rational expectations.

Marcet and Sargent (1995) study the rate at which learning equilibria converge to rational expectations using results from Benveniste et al. (1990), who show that if the real parts of the eigenvalues of the ODE identified by Ljung (1977) are less than $1/2$, then $t^{1/2}(\hat{\theta}_t - \theta)$ has an asymptotic Normal distribution with finite, non-zero variance. However, eigenvalues with real parts greater than $1/2$ can easily arise in practice. We show that the NK model analyzed below has this property in the ZLB. In addition, Marcet and Sargent (1995) use numerical simulations to study versions of the Cagan (1956) model of hyperinflation. In some of those simulations, eigenvalues are substantially larger than $1/2$. Christopeit and Massmann (2018) extend the results in Benveniste et al. (1990) for the Bray and Savin (1986) model to the case of $1/2 \leq b < 1$. We extend some of the results in Christopeit and Massmann (2018) to a class of multivariate models that nests our NK model.

Ferrero (2007) discusses learning in the context of a linear NK model in which the ZLB on interest rates is not binding. He uses the simulation methods proposed by Marcet and Sargent (1995) to study convergence rates of learning equilibria. Ferrero (2007) adopts the so-called Euler-equation approach to learning as opposed to our

approach; see Evans (2021) for a definition of the Euler-equation approach to learning.⁷ Another difference with Ferrero (2007) is that we compare rates of convergence in a nonlinear NK model when the ZLB on interest rates is and is not binding. Ferrero (2007) only considers the latter case.

Cogley and Sargent (2008), Adam and Marcet (2011), and Adam et al. (2017) develop the internalized learning approach in the context of endowment economies. Adam and Merkel (2019) use this approach to analyze a real business cycle model in which peoples’ beliefs do not nest an REE. In contrast, we study an NK model in which peoples’ beliefs do nest an REE and we characterize rates of convergence to that equilibrium.

Preston (2005) and Eusepi et al. (2022) study the effects of monetary policies in linearized NK models under learning, both in and out of the ZLB. They use the anticipated utility approach to modeling how people make decisions. In contrast, we work with a nonlinear model and adopt the internalized learning approach to decision making. Additionally, we characterize how quickly, under learning, monetary and fiscal policies have effects similar to those obtained under rational expectations.

A different literature investigates prices’ information content for fundamentals observed with noise. In this context, Vives (1993) asks a question similar to ours: How quickly do people’s beliefs converge? Specifically, he studies a model in which people use price signals and other noisy observations to learn about an object (a cost parameter) whose value is independent of beliefs. In our model, the values of the objects that people are learning about—for example, aggregate output and inflation—depend on their beliefs. Heemeijer et al. (2009) and Hommes (2011) study positive and negative feedback loops from expectations to outcomes using laboratory experiments and univariate models with constant gain. By contrast, we consider learning in the context of the multivariate NK model under Bayesian learning and analytically characterize the rate of convergence of beliefs.

Our paper is also related to a recent game-theoretic grounded literature that analyzes the implications of bounded rationality for the effectiveness of fiscal and monetary policy. Farhi and Werning (2019) use k -level thinking models to study how deviations from rational expectations affect the effectiveness of forward guidance. García-Schmidt and Woodford (2019) study forward guidance and interest rate pegs using reflective expectations. Iovino and Sergeyev (2023) apply k -level thinking and reflective expect-

⁷Our approach is an example of what Evans (2021) calls the agent-based approach to learning.

tations to analyze the effects of quantitative easing. Angeletos and Lian (2017) develop the idea that a lack of common knowledge can attenuate general-equilibrium effects and damp the effects of government spending. Angeletos and Lian (2017; 2018) analyze the consequences of bounded rationality for the size of fiscal multipliers.

Farhi and Werning (2019), Farhi et al. (2020) and Woodford and Xie (2019; 2022) use different models of bounded rationality to study the size of the government-spending multiplier. Vimercati et al. (2021) assess the implications of bounded rationality for the effectiveness of tax and government spending policy at the ZLB. They do so through the lens of a standard NK model in which people are dynamic k -level thinkers.

In all of the papers just cited, individuals have a limited ability to understand the general equilibrium consequences of monetary and fiscal policies. Like learning, this type of deviation from rational expectations can limit the power of forward guidance. Our paper studies a form of deviation from rational expectations different from those cited in the previous two paragraphs. Moreover, in contrast to our analysis, these papers do not analyze rates of convergence to rational expectations.

3 Simple Example

We consider a workhorse model used in the learning literature; see, for example, Bray and Savin (1986) and Evans and Honkapohja (2001). We use this model to exposit our basic intuition about the factors determining how fast learning models converge. That intuition is summarized by what we refer to as the *learning principle*. First, when people's expectations are partially self-fulfilling then convergence to REE is slow. Second, when people's expectations lead to outcomes that are different from their expectations, then convergence is quick. We exposit this principle using different measures of convergence.

Suppose a variable, x_t , for $t = 1, 2, \dots$, is determined as follows:

$$x_t = a + b\mathbb{E}_{t-1}x_t + \varepsilon_t. \quad (1)$$

Here, ε_t has mean zero, variance $\sigma^2 < \infty$, and is not correlated over time. The operator, $\mathbb{E}_{t-1}$, denotes the cross-sectional average of expectations based on the history of observations on x_t up to period $t - 1$. Evans and Honkapohja (2001) show that when $b > 0$, equation (1) is the reduced form of the Lucas (1973) supply model. When $b < 0$, equation (1) is the reduced form of the Cobweb model analyzed in Muth (1961).

We consider two specifications of $\mathbb{E}_{t-1}$ corresponding to whether people have rational expectations or use past data to learn about the data-generating process for x_t . When $b \neq 0$, the way people form their beliefs affects the law of motion for x_t .

In the REE, $\mathbb{E}_{t-1}$ corresponds to the mathematical expectation, E_{t-1} , and x_t is given by

$$x_t = \mu + \varepsilon_t, \mu = \frac{a}{1-b}. \quad (2)$$

That is, in the REE x_t is distributed iid with mean $a/(1-b)$ and variance σ^2 . For example, if $\varepsilon_t \sim N(0, \sigma^2)$, then $x_t \sim N(a/(1-b), \sigma^2)$.

3.1 Beliefs about the Mean of x_t

As in Bray and Savin (1986) and Evans and Honkapohja (2001), people assume that x_t is Normally distributed with mean μ and variance σ^2 , but they do not know the value of μ .⁸ For ease of exposition, we assume people know the value of σ^2 . The analysis below is unchanged if we assume that people must also learn the value of σ^2 so long as they have Normal-inverse-gamma priors about μ and σ^2 , which would result in the same equations for μ_t .

Another approach to learning used in the related literature is *constant gain learning*. In the model considered here, constant gain learning is not an optimal approach to learning from the perspective of households and firms. As a result, it does not fit into a framework of internalized learning. See Appendix B for further discussion of constant gain learning and for a characterization of the rate of convergence of beliefs to REE.

We assume that before observing x_t , people's prior belief about μ is given by the Normal distribution

$$N(\mu_{t-1}, \sigma^2/\lambda_{t-1}). \quad (3)$$

Here, λ_{t-1} characterizes the precision of the prior about μ .⁹ After seeing x_t people's posterior belief about μ_t is $N(\mu_t, \sigma^2/\lambda_t)$, where

$$\mu_t = \mu_{t-1} + \frac{1}{\lambda_{t-1} + 1} (x_t - \mu_{t-1}), \quad (4)$$

$$\lambda_t = \lambda_{t-1} + 1 = \lambda_0 + t, \quad (5)$$

⁸Recall, we have not assumed that ε_t actually has a Normal distribution. People in the model assume Normality of the likelihood when they derive Bayes' rule.

⁹We can interpret μ_0 as the average of person-specific priors, but, we require that all people have the same value for λ_0 .

for $t = 1, 2, \dots$, where $1/(\lambda_0 + t)$ is the optimal weight on new information.¹⁰ The parameter, λ_0 , is finite and non-negative. If $\lambda_0 = 0$ then μ_t in equation (4) corresponds to the time t least-squares estimator of μ .

Substituting from equation (1), and rearranging, we obtain

$$\mu_t = \frac{a + \varepsilon_t + (b + \lambda_{t-1})\mu_{t-1}}{\lambda_{t-1} + 1}. \quad (6)$$

After repeated substitution, we obtain a decomposition of μ_t in terms of the shocks and μ_0 . Let

$$z_t = \prod_{j=1}^t (1 - b_j), \quad (7)$$

where

$$b_j \equiv \frac{1 - b}{\lambda_0 + j}. \quad (8)$$

The following Lemma summarizes the time-series representation of μ_t that we work with:

Lemma 1. *The variable, μ_t , in equation (6), has the following representation:*

$$\mu_t = \frac{a}{1 - b} + \sum_{j=1}^t \left\{ \frac{z_t}{z_j} \frac{\varepsilon_j}{\lambda_0 + j} \right\} + z_t \left(\mu_0 - \frac{a}{1 - b} \right), \quad (9)$$

where z_t is defined in equation (7).

For the proof, see Appendix A.

3.2 Characterizing Rates of Convergence

Bray and Savin (1986) prove that if $b < 1$, then $\mu_t \rightarrow a/(1 - b)$ almost surely. In contrast, we are interested in the rate at which μ_t converges. To this end, we focus on the mean of the posterior distributions of μ_t . The parameter b is the critical determinant of the rate of convergence.

¹⁰This result about posteriors is well known; see, for example, Hamilton (2020).

3.2.1 Rate of Convergence of Mean of μ_t

We now consider the rate of convergence of $E\mu_t$, where E is the unconditional mathematical expectation. According to Lemma 1,

$$z_t = E \left(\frac{\mu_t - \frac{a}{1-b}}{\mu_0 - \frac{a}{1-b}} \right), \quad (10)$$

for each $t \geq 1$. We can interpret $1 - z_t$ as the fraction of the initial gap, $\mu_0 - a/(1-b)$, closed by period t .

From equations (7) and (8), we see that z_t depends only on λ_0 and b . The smaller the precision, λ_0 , the larger the gain at all dates (see equations (4) and (5)). This *gain effect* implies that the less precise people's initial priors are, the more weight they give to the data and, in our simple model, the more quickly their views converge.

We now analyze numerically how the speed of convergence of $E\mu_t$ depends on b . Our metric is the amount of time it takes to close two-thirds of the initial gap. That is, we calculate T , the value of t such that $z_T \approx 1/3$. In the following calculations, we set $\lambda_0 = 1$.¹¹ When $b = 0, 0.5, 0.75, 0.85$, and $.95$, then $T = 3, 11, 113, 2201$, and 5.2 billion, respectively. Note how the speed of convergence decreases nonlinearly with b . When $b = 0$, people's beliefs converge very quickly. In contrast, when b is large and positive, for example, 0.95 , beliefs essentially take forever to converge.

Suppose $0 < b < 1$. Then, learning injects a *positive feedback loop* into the data. For higher b , a large value of μ_{t-1} implies a large value of x_t (see equation (1)). That, in turn implies a higher value of μ_t (see equation (4)). The higher is b the more powerful is the feedback loop. This positive feedback loop explains why the higher value of b leads to a slower speed of convergence.¹² As we discuss below, this intuition also applies for $b < 0$. That is, for all b satisfying $b < 1$, convergence is faster for smaller values of b .

To discuss the rate of convergence of $E\mu_t$, we need to define what it means for two sequences— x_t and a_t , both of which converge to zero—to have the same rate of convergence. Loosely, two sequences have the same rate of convergence when their ratio does not diverge or converge to zero. This condition is satisfied when (i) the ratio converges to a finite, non-zero constant or (ii) the ratio oscillates in a bounded set. Case (i) is relevant to our analysis of the Bray and Savin (1986) model. As it turns

¹¹Note that the value of a is irrelevant for z_t .

¹²If b is too large, then the feedback loop is too strong, so that the process would not converge. That is the reason why we focus on $b < 1$ in these simulations.

out, case (ii) is relevant to our analysis of the NK model. Our definition accommodates both cases.

Definition 1. Consider two series, x_t and $a_t > 0$ that converge to zero—that is, $\lim_{t \rightarrow \infty} x_t = \lim_{t \rightarrow \infty} a_t = 0$. We say that x_t and a_t converge at the same rate if (a) there exists an $A < \infty$ such that $|x_t|/a_t \leq A$ for all t , and (b) there exists an $\epsilon > 0$ such that for any $T \geq 1$ we have $\sup_{t \geq T} |x_t|/a_t > \epsilon$. If conditions (a) and (b) are satisfied we write $x_t \simeq a_t$.

Conditions (a) and (b) correspond to the requirements that $|x_t|/a_t$ does not diverge and does not converge to zero, respectively.

The following proposition taken from Christopheit and Massmann (2018) establishes the rate at which z_t converges to zero:¹³

Proposition 1. *For any $b < 1$ and any $0 \leq \lambda_0 < \infty$, if $\frac{1-b}{\lambda_0+t} \neq 1$ for all t , then $z_t \simeq t^{b-1}$.*

The requirement that $\frac{1-b}{\lambda_0+t} \neq 1$ for all t is necessary because if $\frac{1-b}{\lambda_0+t^*} = 1$ for some t^* , then $z_t = 0$ for all $t \geq t^*$ (see equation (7)). The assumption that $(1-b)/(\lambda_0+t) \neq 1$ only rules out isolated values of b and λ_0 . While the proof of Proposition 1 is somewhat involved, the intuition can be seen from the following approximations:

$$\log(z_t) \approx -\sum_{j=1}^t b_j \approx c_1 + (b-1) \sum_{j=1}^t \frac{1}{j} \approx c_2 + (b-1) \log(t),$$

where c_1 and c_2 are finite constants. The first approximation is the first-order Taylor series approximation, $\log(1-x) \approx -x$. The second approximation puts the effect of λ_0 in a constant (c_1). The third approximation captures well-known properties of the harmonic sum. If these approximations held with equality then the result of the proposition would follow trivially by exponentiating $\log z_t$. Thus, the power convergence result, that z_t eventually evolves as t^{b-1} , occurs because of the decreasing gain in equation (4). In particular, if the gain were instead constant then z_t would display geometric convergence, *i.e.*, it would eventually evolve according to Λ^t where $|\Lambda| < 1$. It is well known that power convergence is qualitatively slower than geometric convergence.

The analog to the ODE considered in Ljung (1977) that is associated with equation (4) is given by $\dot{\mu}(\tau) = x(\tau) - \mu(\tau)$, where $x(\tau) = b\mu(\tau)$ and τ denotes notional time. The eigenvalue of the mapping from $\mu(\tau)$ to $x(\tau)$ is b . The solution to the ODE is

¹³See equation B.7 of their appendix.

$\mu(\tau) = e^{(b-1)\tau} \mu(0)$. Consistent with Ljung (1977), whether $\mu(\tau)$ converges to zero is determined by the value of b . Proposition 1 shows that b also determines the rate of convergence, in actual time, of $E\mu_t$. Note that the rate of convergence in actual time is a power function of t , where the power is determined by b . This is notable because convergence in notional time is geometric, which is always faster than power convergence.

Proposition 1 says that for large enough t , $|\mu_t - a/(1-b)|$, is approximately κt^{b-1} for some finite constant, $\kappa \neq 0$. So, we can compute how many periods, T_t , it takes to close two-thirds of an initial gap, κt^{b-1} , in period t . It is easily verified that:

$$T_t = \left[3^{\frac{1}{1-b}} - 1 \right] t. \quad (11)$$

Note that T_t only depends on b , and not on other objects like λ_0 . We also compute the time required, T , to close the initial gap in period 0, obtained by simulating the actual z_t 's. When $b = 0, 0.5, 0.75, 0.85$, and 0.95 then $T_1(T) = 2(3), 8(11), 80(113), 1516(2201)$, and 3.5 billion (5.2 billion). It is striking how well T_1 tracks T . We infer that the asymptotic result in Proposition 1 is informative about the behavior of $\mu_t - a/(1-b)$, even for small values of t .

We redo these calculations for the case of $\lambda_0 = 0$ —that is, the case of least-squares learning. For $b = 0, 0.5, 0.75, 0.85$, and 0.95 , we obtain $T_1(T) = 2(1), 8(3), 80(36), 1516(745)$, and 3.5 billion (1.9 billion). When $\lambda_0 = 10$, we obtain $T_1(T) = 2(21), 8(83), 80(831), 1516(15,804)$, and 3.5 billion (36.5 billion). The results for all three values of λ_0 are qualitatively similar. For small values of b , convergence is relatively fast, and for large values of b , the time required to converge explodes.

The following corollary follows immediately from Lemma 1 and Proposition 1.

Corollary 1. *For any $b < 1$ and any $0 \leq \lambda_0 < \infty$: (i) if $\mu_0 \neq a/(1-b)$ and $\frac{1-b}{\lambda_0+t} \neq 1$ for all t , then $E(\mu_t - a/(1-b)) \simeq t^{b-1}$; and (ii) if $\mu_0 = a/(1-b)$ then $E(\mu_t - a/(1-b)) = 0$.*

Corollary 1 establishes the rate of convergence of $E\mu_t$.

3.2.2 Rate of Convergence of Distribution of μ_t

In this paper, we generally focus on convergence in terms of the rate of convergence in mean. Marcet and Sargent (1995) propose an alternative measure of convergence, the rate of convergence in the distribution of μ_t . They characterize that rate by the value

of δ for which $t^\delta (\mu_t - a/(1-b))$ converges to a non-degenerate distribution as $t \rightarrow \infty$. If we knew what that distribution is, then the distribution approach to convergence would provide us with an approximation to the entire distribution of $t^\delta (\mu_t - a/(1-b))$ for small t .

At the time that Marcet and Sargent (1995) was written, there were no analytic results that could be used to develop an approximation for δ , for $b \geq 1/2$. So, Marcet and Sargent (1995) developed a numerical approach. Christopeit and Massmann (2018) develop the required analytic results for the univariate case when $b \geq 1/2$. Their results reveal some limitations of the distribution approach to the rate of convergence. First, when $b < 1/2$, δ is a constant, equal to $1/2$. So, the rate of convergence in distribution does not capture the fact that the rate of convergence of $E\mu_t$ accelerates as b falls below $1/2$. Second, when $b = 1/2$, there is no value δ for which $t^\delta (\mu_t - a/(1-b))$ converges to a non-degenerate distribution as $t \rightarrow \infty$. Third, when $1/2 < b < 1$, $E\mu_t$ converges at the same rate as its standard deviation, implying that the mean of the asymptotic distribution, $\lim_{t \rightarrow \infty} Et^{1-b} (\mu_t - \frac{a}{1-b})$, is not zero. So, the asymptotic distribution may be misleading for small values of t because it is not centered on the point, $a/(1-b)$, to which μ_t converges almost surely. Finally, although Christopeit and Massmann (2018) show that an asymptotic limiting distribution exists when $1/2 < b < 1$, they do not provide a characterization of that distribution. Unless we use numerical simulations, we cannot at this time use the distribution approach to approximate the distribution of $t^\delta (\mu_t - a/(1-b))$ for finite t , when $1/2 < b < 1$.

4 Learning in the New Keynesian Model

In this section, we describe a simple NK model. As in Eggertsson and Woodford (2003), we allow for a shock to the household's discount rate that can cause the ZLB on the interest rate to be binding. To study the properties of the model under learning, it is convenient to express people's problems in recursive form.

In the current period, households discount next period's utility by $1/(1+r)$. In steady state, $r = r_{ss} > 0$. We assume that initially, the economy is in the unique non-stochastic rational expectations steady state in which the nominal interest rate is positive. Then, unexpectedly, $r = r_\ell < r_{ss}$. People correctly understand that next period's discount rate, r' , is drawn from a two-state Markov chain, $r' \in [r_\ell, r_{ss}]$, with

an absorbing state:

$$\begin{aligned}\Pr[r' = r_\ell | r = r_\ell] &= p, \quad \Pr[r' = r_{ss} | r' = r_\ell] = 1 - p, \\ \Pr[r' = r_\ell | r = r_{ss}] &= 0.\end{aligned}\tag{12}$$

Once $r = r_{ss}$, the economy returns to the initial rational expectations steady state. There is another rational expectations steady state in which there is deflation and the nominal interest rate is unity (see Benhabib et al. (2001)). We abstract from that steady state equilibrium because it is not stable under the learning that we consider. Moreover, focusing on one steady state greatly simplifies our analysis. See Arifovic et al. (2018) for a discussion of stability for other learning models in which the deflationary steady state is learnable.

4.1 Fiscal and Monetary Policy

Monetary policy is given by

$$R = \max\{1, 1 + r_{ss} + \alpha(\pi - 1)\},\tag{13}$$

where $\alpha/(1 + r_{ss}) > 1$ and the max operator reflects the ZLB constraint. Later, we discuss other variations on monetary policy including forward guidance.

We consider two specifications for G . In the baseline specification, $G = G_{ss}$, its nonstochastic steady-state value. We also consider a policy where $G = G_\ell > G_{ss}$ while $r = r_\ell$. The government finances its expenditures with lump-sum taxes, $G + \nu wN$, where νwN represents a subsidy paid to intermediate goods firms.

4.2 Private Agents' Problems

Below, we define the household and firm problems.

4.2.1 The Household's Problem When $r = r_\ell$

The household enters a period with a stock of bonds, $b_h = B_{h,t-1}/P_{t-1}$. Here, $B_{h,t-1}$ denotes the beginning-of-period t payoff on nominal bonds acquired in the previous period, when the price of consumption goods was P_{t-1} . At the beginning of a period, before markets open, the household also knows the value the vector, Θ , which

summarizes its beliefs about the distribution of a vector, x :

$$x = \begin{bmatrix} C \\ \pi \end{bmatrix}.$$

Here, C and π denote the current period's aggregate consumption and aggregate inflation. The variable, π , corresponds to P_t/P_{t-1} , where P_t and P_{t-1} denote the current and previous period's aggregate price level, respectively.

In a standard recursive equilibrium, people know current-period market prices and profits when they make their current decisions. Typically, when markets open in these models, people can deduce the prices and profits from a small set of variables. In our context, these variables are the two components of x . In this spirit, we assume that people observe x when markets open, and they make their current consumption, saving, and labor decisions. In making those decisions, households internalize the effect of x on their beliefs about the distribution x' —that is, the value of x in the next period. Those beliefs, Θ' , are given by

$$\Theta' = L(\Theta, x). \quad (14)$$

The form of L depends on the model of learning being analyzed. The household is internally rational in the sense of Adam and Marcet (2011). Specifically, when making decisions, it takes into account uncertainty about the distribution of x and the fact that beliefs about that distribution will evolve as new data arrive (see Section 4.3.2).

Let C_h, N_h, b'_h denote the representative household's consumption, hours worked and end-of-period bond holdings. The household solves

$$\begin{aligned} \max_{C_h, N_h, b'_h} & \left\{ \log(C_h) - \frac{\chi}{2} (N_h)^2 \right. \\ & \left. + \frac{1}{1+r_\ell} [(1-p) V_{h,ss}(b'_h) + p \mathbb{E}_{\Theta'} V_h(b'_h, \Theta', x')] \right\} \end{aligned} \quad (15)$$

subject to

$$C_h + \frac{b'_h}{R(x)} \leq \frac{b_h}{\pi(x)} + w(x) N_h + T(x). \quad (16)$$

Here, $T(x)$ denotes profits net of lump-sum taxes, $w(x)$ denotes the real wage, $R(x)$ denotes the nominal rate of interest, and $\pi(x)$ denotes the inflation rate.¹⁴ In equation

¹⁴We constrain the choice of b'_h to a compact set $[\underline{b}, \bar{b}]$, which we discuss in Appendix C.

(15), $V_{h,ss}(b'_h)$ denotes the value function of the household conditional on $r' = r_{ss}$ and $V_h(b'_h, \Theta', x')$ denotes the value conditional on $r' = r_\ell$. The expectation operator, $\mathbb{E}_{\Theta'}$, is evaluated using the marginal data density for x' implied by $\Theta' = L(\Theta, x)$ and $r' = r_\ell$. Using the first order optimality condition for N_h and equation (16), we reduce the household problem to finding an optimal decision rule, $b'_h(b_h, \Theta, x)$, for bond holdings.

The function, $V_{h,ss}(b_h)$, satisfies the following fixed point:

$$V_{h,ss}(b_h) = \max_{C_h, N_h, b'_h} \left\{ \log(C_h) - \frac{\chi}{2} (N_h)^2 + \frac{1}{1 + r_{ss}} V_{h,ss}(b'_h) \right\}, \quad (17)$$

subject to

$$C_h + \frac{b'_h}{R_{ss}} \leq \frac{b_h}{\pi_{ss}} + w_{ss} N_h + T_{ss},$$

where T_{ss} denotes steady-state profits net of taxes in steady state, w_{ss} denotes the steady-state real wage, R_{ss} denotes the steady-state nominal interest rate, and π_{ss} denotes the steady-state inflation rate.

The function, V_h , in equation (15) has the fixed point property:

$$V_h(b_h, \Theta, x) = \max_{C_h, N_h, b'_h} \left\{ \log(C_h) - \frac{\chi}{2} (N_h)^2 + \frac{1}{1 + r_\ell} [(1 - p) V_{h,ss}(b'_h) + p \mathbb{E}_{\Theta'} V_h(b'_h, \Theta', x')] \right\}, \quad (18)$$

where the maximization is subject to equation (16) and the law of motion for Θ in equation (14).

4.2.2 The Firm's Problem When $r = r_\ell$

A final homogeneous good, Y , is produced by competitive and identical firms using the technology

$$Y = \left(\int_0^1 Y_f^{\frac{\varepsilon-1}{\varepsilon}} df \right)^{\frac{\varepsilon}{\varepsilon-1}}, \quad (19)$$

where $\varepsilon > 1$. The representative firm chooses inputs, Y_f , to maximize profits $YP - \int_0^1 Y_f P_f df$, subject to (19). The firm's first order condition for the f^{th} input is

$$Y_f = \left(\frac{P_f}{P} \right)^{-\varepsilon} Y. \quad (20)$$

The f^{th} intermediate good is produced by a monopolist with production technology $Y_f = N_f$, where N_f is labor hired by firm f . Let p_f denote the f^{th} firm's price in the previous period, scaled by that period's aggregate price index—that is, $P_{f,t-1}/P_{t-1}$. Also, let p'_f denote the firm's current choice of price scaled by the current aggregate price index. In our scaled notation,

$$\frac{p'_f}{p_f} \pi = \frac{P_{f,t}}{P_{f,t-1}}. \quad (21)$$

Firms value a unit of real profits by the marginal utility of consumption, $1/C$. Prices are sticky as in Rotemberg (1982). When $r = r_\ell$ the current-period problem of firm f is to set its price p'_f so that

$$\begin{aligned} p'_f(p_f, \Theta, x) = \operatorname{argmax}_{p'_f} & \frac{1}{C(x)} \left\{ (p'_f - (1 - \nu) w(x)) (p'_f)^{-\varepsilon} Y(x) \right. \\ & \left. - \frac{\phi}{2} \left(\frac{p'_f}{p_f} \pi(x) - 1 \right)^2 (C(x) + G(r_\ell)) \right\} \\ & + \frac{1}{1 + r_\ell} [(1 - p) V_{f,ss}(p'_f) + p \mathbb{E}_{\Theta'} V_f(p'_f, \Theta', x')]. \end{aligned} \quad (22)$$

Here, $V_{f,ss}(p'_f)$ denotes the value of the firm's problem conditional on $r' = r_{ss}$ and $V_{f,ss}(p'_f, \Theta', x')$ denotes its value conditional on $r' = r_\ell$.¹⁵ Firms and households have the same information sets and update priors in the same way. Thus, the expectations operator is the same as the one in the household's problem. In equation (22), we follow the literature by scaling price adjustment costs by real GDP.¹⁶ Also, ν is a tax subsidy on employment designed to eliminate the effect of monopoly distortions in steady state.¹⁷

The function, $V_{f,ss}(p_f)$, has the fixed-point property

$$\begin{aligned} V_{f,ss}(p_f) = \max_{p'_f} & \left\{ \frac{1}{C_{ss}} \left((p'_f - (1 - \nu) w_{ss}) (p'_f)^{-\varepsilon} Y(x) \right) \right. \\ & \left. - \frac{1}{C_{ss}} \frac{\phi}{2} \left(\frac{p'_f}{p_f} \pi_{ss} - 1 \right)^2 (C_{ss} + G_{ss}) + \frac{1}{1 + r_{ss}} V_{f,ss}(p'_f) \right\}. \end{aligned} \quad (23)$$

¹⁵We constrain the choice of $\log(p'_f)$ to a compact set $[\underline{p}, \bar{p}]$. See Appendix C for a discussion.

¹⁶See, for example, Kaplan and Violante (2018, page 711).

¹⁷That is, $(1 - \nu)\varepsilon/(\varepsilon - 1) = 1$.

The function, V_f , in equation (22) has the fixed point property

$$V_f(p_f, \Theta, x) = \max_{p'_f} \left\{ \frac{1}{C(x)} (p'_f - s) (p'_f)^{-\varepsilon} Y(x) - \frac{1}{C(x)} \frac{\phi}{2} \left(\frac{p'_f}{p_f} \pi(x) - 1 \right)^2 (C(x) + G(r_\ell)) + \frac{1}{1+r_\ell} [(1-p) V_{f,ss}(p'_f) + p \mathbb{E}_{\Theta'} V_f(p'_f, \Theta', x')] \right\}. \quad (24)$$

The maximization takes into account the law of motion of Θ , L , in equation (14).

4.2.3 The Mapping from x to Aggregate Variables

For individual households' and firms' problems to be well defined, they must know the values of seven aggregate variables, $\begin{bmatrix} C & \pi & R & Y & N & w & T \end{bmatrix}$. We assume that each agent knows the model's static equilibrium conditions so they can deduce those variables from $x = \begin{bmatrix} C & \pi \end{bmatrix}$. We denote this mapping by $F(x)$. Households derive R from π using equation (13). The mappings from x and r to Y , N , and w are given by

$$Y = (C + G(r)) \left(1 + \frac{\phi}{2} (\pi - 1)^2 \right), \quad N = Y, \quad w = \chi NC.$$

The first two equalities correspond to goods market clearing and the aggregate production function. The third equality corresponds to the belief that the labor supply curve of the individual household holds as an aggregate condition. These equalities hold in every period of our learning equilibria (described in the next sub-section).

Aggregate firm profits net of taxes implied by x and r are

$$T = (1 - w) Y - \frac{\phi}{2} (\pi - 1)^2 (C + G(r)) - G(r).$$

4.3 Equilibrium and Beliefs

The equilibrium for our model is a *learning equilibrium* for the duration of time that $r = r_\ell$, followed by a jump to the positive interest rate, steady state REE. The learning equilibrium is a sequence of period equilibria.

4.3.1 Equilibrium Definitions

We now define a temporary equilibrium.

Definition 2. Given Θ and r_ℓ , a period equilibrium is a set of values of x and $\Theta' = L(\Theta, x)$ such that

- (i) households and firms solve their optimization problems, defined in equations (15) and (22), respectively
- (ii) labor, goods and bond markets clear
- (iii) $p'_f = 1$, $C_h = C$, $N_h = N$

Because firms are identical, in a learning equilibrium, no firm will ever inherit a $p_f \neq 1$. Then, equation (21) and the first part of condition (iii) imply that people's views about inflation, π , are correct. The second and third parts of condition (iii) imply that people's views about C and N are correct.

The only new conditions in Definition 2 relative to those imposed by $F(x)$ are that bond markets clear ($b'_h = 0$) and firms choose $p'_f = 1$. These two conditions determine the two elements of x .

Two comments about the period equilibrium are worth emphasizing. First, people have perfect foresight regarding current aggregate variables. Second, in general, they do not have perfect foresight about future aggregates. It follows that the period equilibrium under learning is, in general, different from what it would be if people had rational expectations.

We now define a learning equilibrium.

Definition 3. A *learning equilibrium* is :

- (i) a sequence of period equilibria in which beliefs are updated according to equation (14) when $r = r_\ell$,
- (ii) a steady state REE with $R > 1$, when $r = r_{ss}$.

In a learning equilibrium, the value of Θ in the first period when $r = r_\ell$ is exogenous. We assume that in the case of an unprecedented event, people's priors about the economic variables, x , are very diffuse. Below, we describe how our parameterization of the initial Θ captures this property.

4.3.2 Beliefs and Equilibrium

We now describe in detail how households' and firms' common beliefs evolve, starting in the first period that $r = r_\ell$. People assume that each of the two elements of $\log(x)$

is drawn from a Normal distribution:

$$\log(x) = \begin{bmatrix} \log(C) \\ \log(\pi) \end{bmatrix} = \begin{bmatrix} \mu_C \\ \mu_\pi \end{bmatrix} + \begin{bmatrix} \varepsilon_C \\ \varepsilon_\pi \end{bmatrix}, \quad (25)$$

$E\varepsilon_C = E\varepsilon_\pi = 0$, $E\varepsilon_C^2 = \sigma_C^2$ and $E\varepsilon_\pi^2 = \sigma_\pi^2$. These distributions are independent across time and the elements of $\log(x)$. People are uncertain about the values of μ_i , σ_i^2 for $i \in \{C, \pi\}$. Their prior about μ_i conditional on σ_i^2 is Normal, parameterized with a mean, m_i , and variance, σ_i^2/λ_i , where λ_i characterizes the precision of the prior about μ_i . The marginal density of their prior for σ_i^2 is proportional to an inverse-gamma distribution, with shape and scale parameters, α_i and $(\psi_i^2(\alpha_i + 1/2))$, respectively. The prior for σ_i^2 is not exactly an inverse-gamma distribution because we truncate the support of σ_i^2 so that $E[C]$ and $E[\pi]$ have finite values. We find it convenient to express the scale parameter in this way because ψ_i is a consistent estimator for σ_i . The joint density of μ_i, σ_i^2 is proportional to the Normal inverse-gamma distribution. We collect the parameters of the priors in the vector Θ :

$$\Theta = \left(m_C \quad m_\pi \quad 1/\lambda_C \quad 1/\lambda_\pi \quad \psi_C \quad \psi_\pi \quad 1/\alpha_C \quad 1/\alpha_\pi \right). \quad (26)$$

The posterior distribution is also proportional to the Normal inverse-gamma distribution, and the function, L , in equation (14) can be constructed using standard updating formulas, which are detailed in Appendix C.

4.3.3 Anticipated Utility

Virtually all of the related literature works with a version of Kreps' *Anticipated Utility* approach to how people integrate learning into their decisions. While this approach has computational advantages, it has been criticized for being internally inconsistent (see Cogley and Sargent (2008) and Adam and Marcet (2011)). We assess the robustness of our results to using the anticipated utility approach. In our context, that approach assumes that when households and firms make their state- x contingent decisions, they assume that in the current and all future periods, $\log(x)$ will be drawn from a Normal distribution with mean and variance *fixed* at the values of m_i and ψ_i^2 from the beginning-of-period Θ . We make two changes to the household and firm decision problems to implement this assumption. First, we set $\Theta' = \Theta$ in their next-period value functions. Second, in evaluating the expectation operator, $\mathbb{E}_{\Theta'}$, that appears in

the household and firm problems, we use the log Normal density for x with mean and variance fixed at the values of m_i and ψ_i^2 from Θ . Importantly, at the beginning of the next period, firms and households set $\Theta' = L(\Theta, x)$.

In sum, anticipated utility differs from internalized learning in two ways. First, in making their state- x contingent decisions, people ignore the fact that after they see current x , they will update their views, using $\Theta' = L(\Theta, x)$. Second, they ignore their uncertainty about the mean and variance of the distribution of $\log(x)$.

5 Multiple Rational Expectations Equilibria

In this section, we describe the equilibria in our model when agents have rational expectations.

An equilibrium is a set of values for output, employment, inflation, and consumption, $Y_\ell, N_\ell, \pi_\ell, C_\ell$, respectively, when $r = r_\ell$. We assume that the economy reverts to the unique rational equilibrium steady state, $Y_{ss}, N_{ss}, \pi_{ss}, C_{ss}$, with $R_{ss} > 1$ when $r = r_{ss}$.¹⁸

The four equilibrium conditions associated with the four unknowns, $\pi_\ell, C_\ell, R_\ell, N_\ell$, are

$$1 = \frac{1}{1 + r_\ell} \left[p \frac{1}{\pi_\ell} + (1 - p) \frac{C_\ell}{C_{ss}} \right], \quad (27)$$

$$(\pi_\ell - 1) \pi_\ell (C_\ell + G_\ell) = \frac{\varepsilon - 1}{\phi} (\chi N_\ell C_\ell - 1) N_\ell + \frac{1}{1 + r_\ell} p (\pi_\ell - 1) \pi_\ell (C_\ell + G_\ell), \quad (28)$$

$$N_\ell = (C_\ell + G_\ell) \left(1 + \frac{\phi}{2} (\pi_\ell - 1)^2 \right), \text{ and} \quad (29)$$

$$R_\ell = \max \{1, 1 + r_{ss} + \alpha (\pi_\ell - 1)\}. \quad (30)$$

In equations (27) and (28) we have taken into account that $\pi_{ss} = 1$. In addition, we verify and use the fact that $R_\ell = 1$. Equation (27) can be expressed as one equation in the unknown, π_ℓ , after using equations (27) and (29), to express C_ℓ and N_ℓ as functions of π_ℓ . We compute C_{ss} using the steady state of the model. Then, we can find a candidate equilibrium by finding a value of π_ℓ that sets a function, $f(\pi_\ell) = 0$.

¹⁸Throughout the paper, we only consider equilibria in which quantities and prices are constant for a given value of r . For example, we do not consider sunspot equilibria.

To verify that a candidate value of π_ℓ is an equilibrium, we must verify that the implied aggregate quantities and firm values are non-negative.

Our baseline parameters are:

$$p = 0.80, r_\ell = -0.0015, G_{ss} = G(r_{ss}) = 0.20, \beta = 0.995,$$

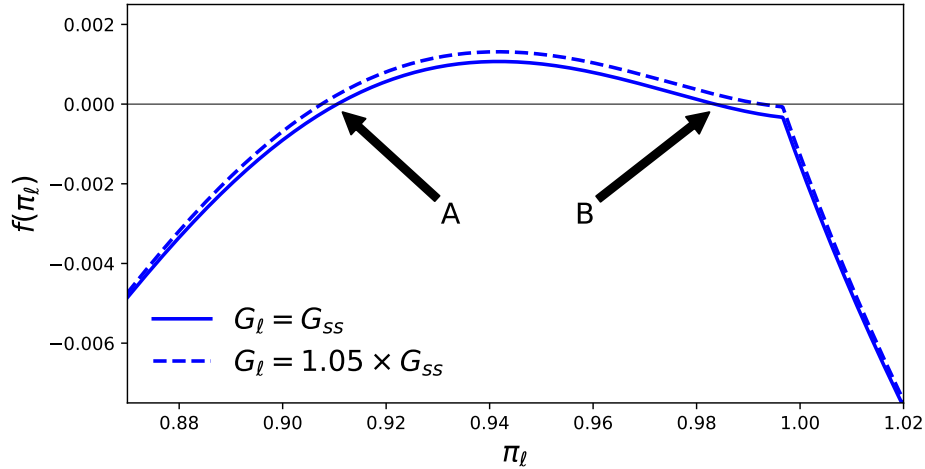
$$\varepsilon = 4, \phi = 110, \chi = 1.25, \alpha = 1.5$$

In the $R > 1$ steady-state REE, $C_{ss} = 0.8, \pi_{ss} = 1, N_{ss} = 1$. In the alternative specification of government purchases,

$$G_\ell = G(r_\ell) = 1.05 \times G(r_{ss}), \quad (31)$$

while $r = r_\ell$.

Figure 1: $f(\pi)$ Corresponding to the Target-Inflation Steady-State Equilibrium



Note: The function, f , is defined in the text. The dashed line is discussed in Section 8.1 below. The range of π_ℓ in the figure includes the two values of π_ℓ that correspond to an equilibrium. Source: Authors' calculations.

Figure (1) displays the function $f(\pi_\ell)$ for a range of values of π_ℓ in the baseline (solid blue line) and alternative (dashed blue line) cases. In each case, there are two values of π_ℓ for which $f(\pi_\ell) = 0$. Table 1 reports the values of $C_\ell, w_\ell, N_\ell, R_\ell$ and π_ℓ at these zeros of f . Each crossing corresponds to an interior equilibrium in which the ZLB binds. The values of the variables corresponding to the equilibria in Figure (1) are reported in Table 1.

The economy's response to a drop in r is the result of two countervailing forces.

Table 1: Equilibrium Values While $r_t = r_\ell$, Returning to Target-Inflation Steady State

Label	Bad ZLB	Good ZLB
	A	B
$400(\pi^\ell - 1)$	-35.78	-6.60
$400(R^\ell - 1)$	0	0
C^ℓ	0.48	0.74
N^ℓ	0.98	0.95
w^ℓ	0.59	0.88

(a) $G_\ell = G_{ss}$

Label	Bad ZLB	Good ZLB
	A	B
$400(\pi^\ell - 1)$	-36.99	-3.00
$400(R^\ell - 1)$	0	0
C^ℓ	0.47	0.77
N^ℓ	1.00	0.98
w^ℓ	0.58	0.95
$\frac{\Delta C + \Delta G}{\Delta G}$	-0.17	3.95

(b) $G = 1.05 \times G_{ss}$

Note: This table reports $\{\pi_\ell, R_\ell, C_\ell, N_\ell, w_\ell\}$ for two equilibria indicated by A and B when $G = G_{ss}$ (2a) and when $G = 1.05G_{ss}$ (2b). Each equilibrium returns to the target-inflation steady state as soon as $r = r_{ss}$. The government purchases multiplier reported in the last line of panel is the change in GDP per unit increase in G within each of the type A and B equilibria. Source: Authors' calculations.

First, the drop in r leads to an increase in desired savings. In the first best equilibrium, the real interest rate would drop enough to undo the increased desire to save completely, allowing market clearing in the bond and goods market without any change in consumption and employment. When monetary policy is operated by a Taylor rule, and prices are sticky, then we know that policy goes only part-way towards achieving the first best equilibrium. The real interest rate falls, but not by enough so that market clearing must be accomplished in part by a drop in output and income, which reduces the desire to save, as long as the low- r spell is expected to be short enough (that is, p is small enough).¹⁹ If the required fall in the nominal interest rate is sufficiently large, then the ZLB on the nominal interest rate binds. When the ZLB binds, a form of *deflation spiral* is triggered. The fall in output leads to a drop in marginal cost that reduces actual and expected deflation. The latter raises the real interest rate, amplifying the desire to save, leading to an additional drop in actual and expected inflation.

¹⁹Further discussion of this point appears below.

An important countervailing force limits the extent of this spiral. As output drops, consumption smoothing leads people to save less. The lower is p , the shorter is the expected duration of the ZLB and the stronger is the consumption smoothing motive.

Three observations about the ZLB follow. First, the logic of the deflation spiral provides intuition into why the fall in output can be very large when the ZLB is binding. The larger the expected deflation in an REE, the larger is the drop in output. Second, the interplay between the deflationary spiral and consumption smoothing provides intuition for why there can be multiple REEs in the ZLB. Third, if p is sufficiently large, the consumption smoothing motive is very weak. When the deflationary spiral is too dominant, an REE does not exist.²⁰

Turning to the fiscal multiplier, we calculate the effect of an increase in G comparing A to A' and B to B' —that is, comparing two Bad-ZLB equilibria and two Good-ZLB equilibria (see Figure 1). Table 1 shows that the multiplier is very large in the latter case and very small in the former. Consistent with this observation, expected deflation is much larger at A' than at B' .

In sum, this section highlights the central role that expected deflation plays in determining the properties of an REE in the ZLB. We expect that because expectations are backward looking, the properties of the learning equilibrium will be very different from those of the REE.

6 Equilibrium Selection

In this section, we consider whether the multiplicity of REEs can be resolved by learnability. We analyze the learnability of an REE by considering a small perturbation in the REE beliefs. We consider these perturbations by analyzing learning equilibria with initial values of Θ that are not REE beliefs, but are close to the REE values. We say that an REE is *learnable* if learning equilibria that begin with beliefs in a neighborhood of the REE beliefs converge to the REE. In this section, we conduct the analysis numerically and consider initial values for Θ that have m_C and m_π equal to $\log(C_\ell)$ and $\log(\pi_\ell)$, respectively, where C_ℓ and π_ℓ are the REE values of C and π . Importantly, the variance of the priors is greater than zero. If learning equilibria starting with these values of Θ converge to the associated REE, then we say that the REE is learnable. We have also considered learning equilibria that begin with a vector Θ in which m_C

²⁰See Werning (2012), who also discusses the possibility of nonexistence of equilibrium in the ZLB.

and m_π are near, but not equal to, the associated REE values. In these cases, we find similar results and our conclusions about learnability are unchanged.

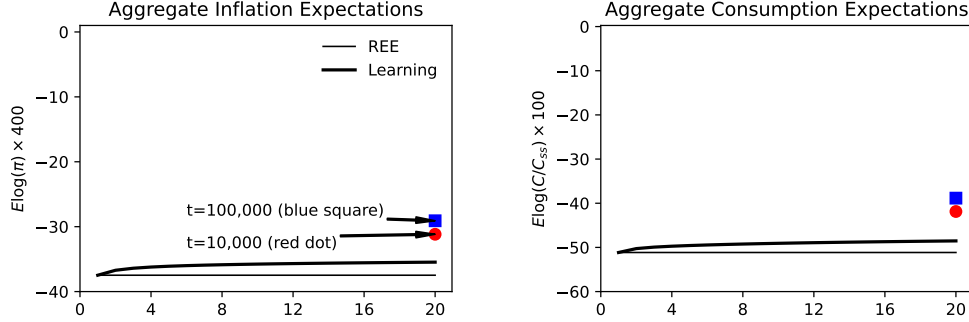
Other initial values of Θ are of particular interest. For example, beliefs with m_C and m_π equal to $\log(C_{ss})$ and $\log(\pi_{ss})$, respectively, are natural candidates in the initial values of Θ . If an REE is learnable and learning equilibria beginning with these initial values of Θ also converge to that REE, then we say that the REE is quasi-globally learnable. In a model with multiple REEs (like the NK model), any particular REE cannot be globally learnable. This result obtains because if beliefs are consistent with another REE, then beliefs will not diverge from that equilibrium.

We initially consider the learnability of the Bad-ZLB equilibrium by examining a learning equilibrium with m_i set to the Bad-ZLB equilibrium values. Figure 2a suggests that the learning equilibrium deviates from the Bad-ZLB equilibrium. The red dot shows where that equilibrium is after 10,000 periods and indicates that it is headed toward the Good-ZLB equilibrium. In Section 9, we use linearization methods to prove that at the assumed parameter values, the learning equilibrium cannot converge to the Bad-ZLB equilibrium.²¹ We conclude that the Bad-ZLB equilibrium is not learnable.

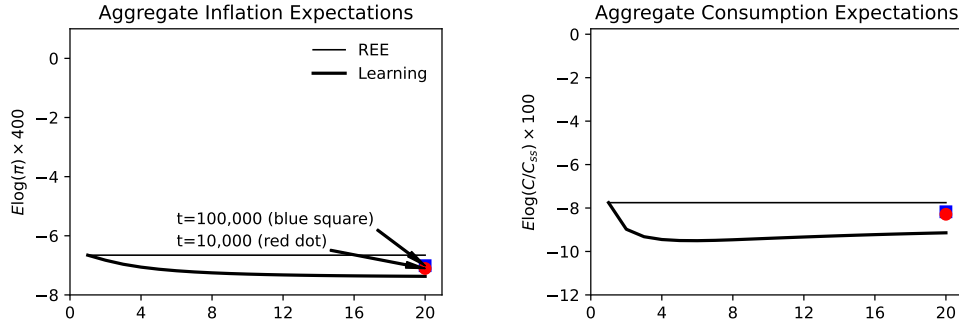
²¹Our proof is by contradiction. We linearize our learning model around the Bad-ZLB equilibrium. Suppose the Bad-ZLB equilibrium is stable. Then, the learning equilibrium would eventually (as long as $r = r^\ell$) arrive in an arbitrary small interval, U , about the Bad-ZLB equilibrium, where our linearized system is arbitrarily accurate. We show that that model satisfies the conditions of Theorem 7.2 in Evans and Honkapohja (2001) for beliefs to leave U . This outcome contradicts the hypothesis that the Bad-ZLB equilibrium is stable.

Figure 2: Equilibrium Selection in the ZLB, by Learning

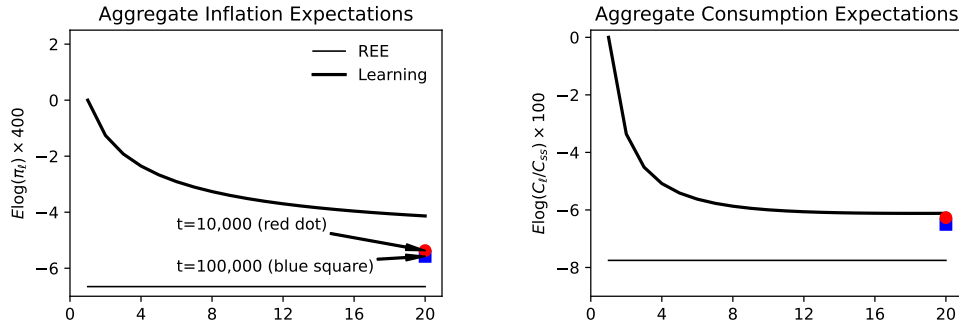
(a) Non-learnability of Bad-ZLB Equilibrium



(b) Learnability of Good-ZLB Equilibrium



(c) Learnability of Good ZLB Equilibrium



Note: In the panels (a) and (b), m_i is initially set to the associated REE value. In panel (c) m_i is initially set to the steady state REE value. In all sub-figures, $\psi_i = 0.02$, $\lambda_i = 1$, $\alpha_i = 2$. Source: Authors' calculations.

Figure 2b shows that the learning equilibrium is converging to the Good-ZLB equilibrium. In Section 9, we use linearization methods to prove that at the assumed parameter values, the learning equilibrium will converge to the Good-ZLB equilibrium if beliefs start in a neighborhood of that REE. Figure 2c shows that the learning equi-

librium converges to the Good-ZLB equilibrium when the beliefs are initially centered on the steady-state REE. Taken together, these results indicate that the Good-ZLB is quasi-globally learnable.

7 Speed of Convergence

In this section, we consider how quickly the learning equilibrium converges to the unique learnable REE. In the first subsection, we consider our results for the baseline parameterization of the model. In the second subsection, we consider the effect of the ZLB on the interest rate on the speed of learning.

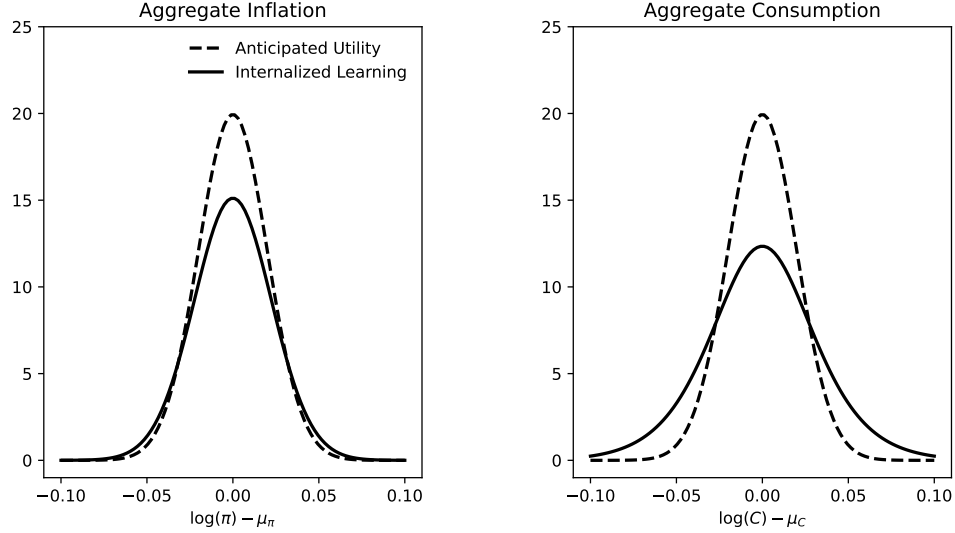
7.1 Baseline Results

We now consider the effects of a drop in r in our learning model. Our basic assumption is that when people are confronted with an unprecedented observation, here modeled as a drop in r , they become very uncertain about how market-determined variables will evolve. We set the initial value of Θ to the following vector:

$$\begin{aligned} & \left(m_C \quad m_\pi \quad 1/\lambda_C \quad 1/\lambda_\pi \quad \psi_C \quad \psi_\pi \quad 1/\alpha_C \quad 1/\alpha_\pi \right)' \\ &= \left(\log(C_{ss}), \quad \log(\pi_{ss}), \quad 1, \quad 1, \quad 0.02, \quad 0.02, \quad 1/2, \quad 1/2 \right)'. \end{aligned}$$

Figure 3 displays the marginal density of $\log C$ and $\log \pi$ associated with anticipated utility (that is, the Normal distribution evaluated at the prior estimates of the means and variances) and with internalized learning (that is, the marginal data density associated with the Normal-inverse-gamma prior on the parameters of the Normal distribution). Note the fatter tails on the density function associated with internalized learning. The tails are fatter for consumption than inflation because we set a higher upper bound on σ_C (0.05) than on σ_π (0.025). The bounds on the standard deviations correspond to typical period-by-period shock sizes equal to about 6 percent for aggregate consumption and about 10 percentage points for *annualized* aggregate inflation.

Figure 3: Data Density Under Two Models of the Interaction of Beliefs and Decisions



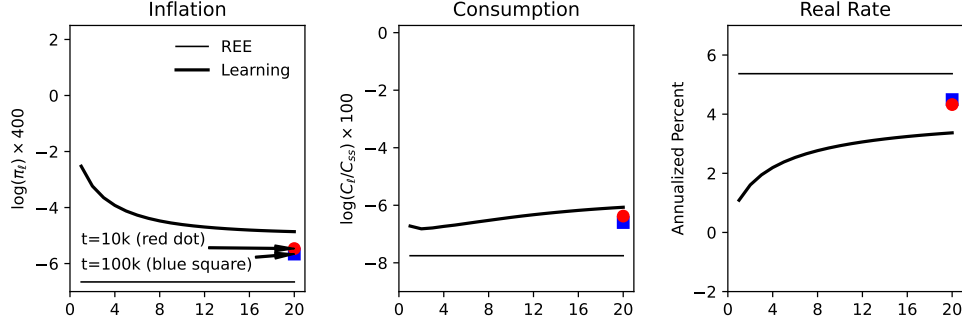
Note: The dashed line corresponds to Normal density functions with means m_i and standard deviations ψ_i . The solid line corresponds to the marginal data density of $\log(x)$ at time one, using Θ before it is updated by the time one value of x is realized. Source: Authors' calculations.

The thin and thick solid lines in Figure 4 display the evolution of inflation, consumption, and the real interest rate after the drop in r under REE and learning, respectively. First, consider Figure 4a, which reports results for the REE and internalized learning. Two key features are worth noting. First, in the REE, there is a very large drop in inflation and consumption, and the real interest rate rises sharply. The fall in inflation and consumption and the rise in the real rate are much smaller under learning. Second, the learning economy converges very slowly to the REE. As shown in Figure 2c, after people initially change their views somewhat quickly, the rate at which they change their views slows dramatically. For example, the dot labeled $T = 10,000$ displays people's views about the variables after 10,000 quarters. Given our value, $p = 0.8$, r is only expected to be low for about five quarters. Whether convergence to the REE happens after 20 quarters or 10,000 quarters is irrelevant. The crucial point is that in a typical ZLB episode, people's beliefs are very far from rational expectations.

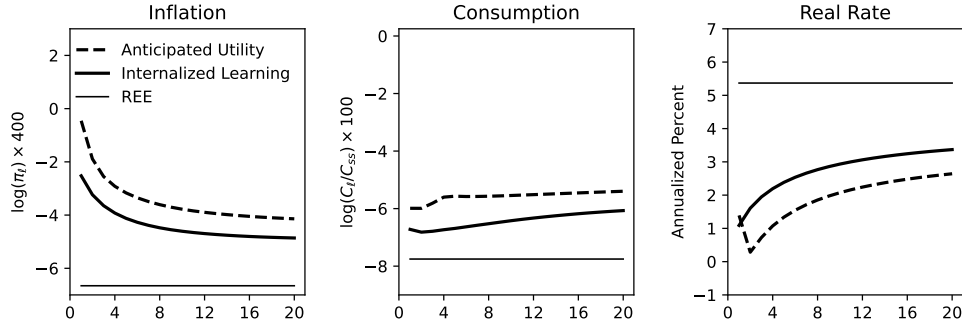
Now consider Figure 4b. This figure compares the evolution of the learning equilibrium under anticipated utility (dashed line) and internalized learning (solid line). The key takeaway is that we obtain the same slow-learning result qualitatively regardless of which approach we take to learning. However, consumption and inflation fall somewhat more under internalized learning.

Figure 4: Simulations of Benchmark Model

(a) Speed of Convergence in the Benchmark Model



(b) Comparison, Speed of Convergence Under Anticipated Utility and Benchmark

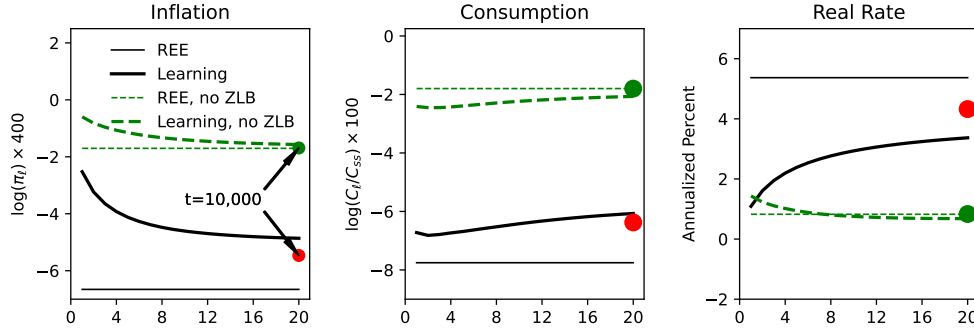


Source: Authors' calculations.

7.2 The Role of the ZLB in the Baseline Results

Figure 5 reports a simulation of our benchmark model in which the ZLB on the interest rate is ignored. For convenience, we reproduce the results from Figure 5 in which the ZLB is binding. The key result is that the learning economy converges very quickly when the ZLB is not binding. The reason is that the Taylor rule weakens the connection between expected and realized inflation. To understand why, suppose people's prior is that inflation will be high in the next period, causing firms to want to raise prices in the current period. When the Taylor principle is operative, the central bank takes actions in the current period that make actual inflation lower. Because expectations are less self fulfilling, the learning principle implies that people will quickly adjust their beliefs. The speed with which they do so depends very much on the value of α , a point that we return to in Section 9.

Figure 5: Benchmark Simulations with and without Binding ZLB



Source: Authors' calculations.

8 Learning and Government Policy

In this section, we analyze the sensitivity of monetary and fiscal policy analysis in the ZLB to deviations from rational expectations. We juxtapose that sensitivity to the lack of sensitivity when the ZLB is not binding.

8.1 Government Purchases Multiplier

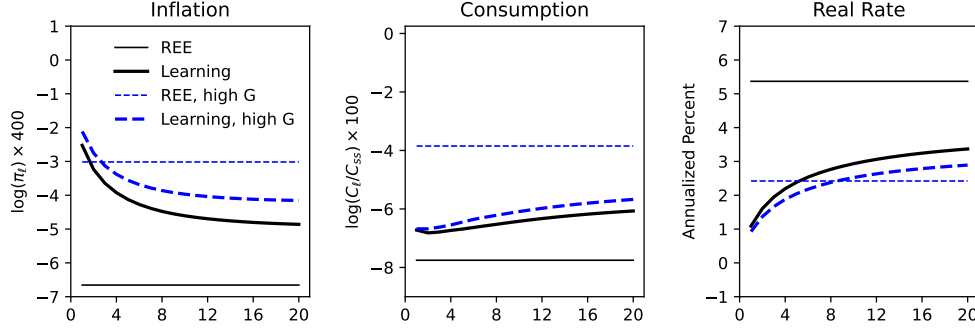
We begin by analyzing the effect of learning on the government purchases multiplier when the ZLB binds. We compute the multiplier by considering the effect on GDP, $C+G$, of a 5 percent rise in government purchases relative to its steady-state level—that is, $G(r_\ell) = 1.05 \times G(r_{ss})$. We denote the difference in consumption and government purchases across the two equilibria by ΔC and $\Delta G = 0.05 \times G(r_{ss})$. Since the Bad-ZLB equilibrium is not stable under learning, we focus on ΔC across Good-ZLB equilibria. We define the multiplier as

$$\frac{\Delta C + \Delta G}{\Delta G}. \quad (32)$$

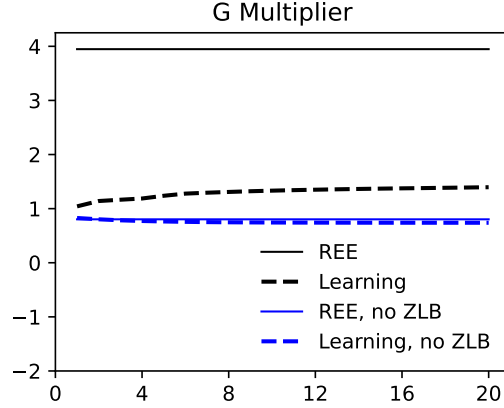
In the REE, the multiplier in the ZLB is equal to 3.95 (see Figure 6b). The multiplier is large when the ZLB is binding because the rise in G generates an increase in expected inflation (see the left panel in Figure 6b). Because R is fixed, this rise generates a fall in the real interest rate and a rise in C (see the middle panel). So, in this case, the multiplier is bigger than one.

Figure 6: Equilibria with and without Jump in G

(a) Increase in Government Purchases During ZLB



(b) Government Purchases Multiplier in ZLB



Notes: The solid lines in Figure 6a reproduce the results based on $G = G_{ss}$ in Figure 4a. The dashed lines report the simulation of the model when $G = 1.05 \times G_{ss}$. Figure 6b displays the government purchases multiplier under internalized learning and in the REE. That figure reports results for the case in which the ZLB is imposed and not imposed ('no ZLB').

Under learning, expected inflation is backward-looking and does not move much with a rise in G (compare the thick dashed and thick solid lines in the left panel of Figure 6a). Hence, the real interest rate does not fall very much and the response in consumption is small (middle panel).

Figure 6b displays the value of the multiplier over time in the REE and under learning in the ZLB. Consistent with the results above, the multiplier in the learning case is small compared with what it is in the REE. Significantly, there is very little convergence of the learning multiplier to its REE value over the 20 quarters displayed.

We now turn to the case when the ZLB is not binding. The REE multiplier, in this case, is much smaller, 0.80, than when the ZLB is binding (see Figure 6b). When the

ZLB is not binding, the rise in inflation causes the monetary authority to raise the real interest rate, which leads to a fall in C . That rise is why the REE multiplier is less than unity outside the ZLB. Figure 6b displays the government purchases multiplier in the learning equilibrium when the ZLB is ignored. Note that the value of the multiplier is very similar to its value in the REE. This result is not surprising in light of our demonstration that when the ZLB is not binding, the learning equilibrium converges quite quickly to the REE.

8.2 Forward Guidance

In this subsection, we consider the sensitivity of the effects of forward guidance to learning. Under such a policy, the monetary authority commits to keeping the nominal interest rate at the ZLB for J periods after the discount rate has returned to its steady-state level. To make our point as simply as possible, we consider the case $J = 1$. In the first subsection we show that the number of REE proliferates under forward guidance. Only one of those equilibria is stable under learning. Second, we analyze the effect of forward guidance.

8.2.1 Rational Expectations Equilibria

We construct the REEs with forward guidance by working backward in three steps. First, we compute the unique non-stochastic steady state with $R > 1$. Second, we compute the continuation equilibrium in the period, I , in which r switches from r_ℓ to r_{ss} , where $I \in [2, 3, \dots]$. Third, we compute the equilibrium allocations in the periods before I , denote by I_{-1} .

In period I , $R = 1$ even though $r = r_{ss}$. People know that the economy will transition to steady state in period $I + 1$. The equilibrium conditions in period I are equations (27) through (29) adjusted for forward guidance:

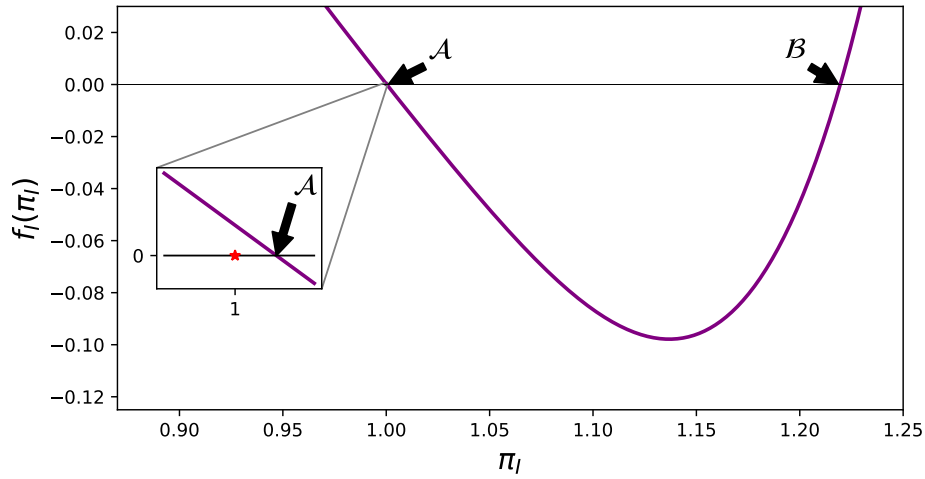
$$1 = \frac{1}{1 + r_{ss}} \frac{C_I}{\pi_{ss} C_{ss}} \quad (33)$$

$$(\pi_I - 1) \pi_I (C_I + G_{ss}) - \frac{\varepsilon - 1}{\phi} (\chi N_I C_I - 1) N_I = 0 \quad (34)$$

$$N_I = (C_I + G_{ss}) \left(1 + \frac{\phi}{2} (\pi_I - 1)^2 \right). \quad (35)$$

Equations (33) and (35) define functions mapping π_I to C_I and N_I . These functions allow us to express the left-hand side of equation (34) as a function of π_I . We denote this function by $f_I(\pi_I)$. A candidate continuation equilibrium in period I is a value of π_I such that $f_I(\pi_I) = 0$ along with the associated values of C_I, N_I, w_I and the present value of the intermediate good firm in period I . For a candidate equilibrium to be an equilibrium, the four variables must be non-negative. Figure 7 displays the f_I function for a range of values of π_I . We find two continuation equilibria corresponding to the two zeros of f_I displayed in the figure (see points $\mathcal{A}$ and $\mathcal{B}$).²²

Figure 7: Equilibria in Period of Switch from $r = r_\ell$ to $r = r_{ss}$ Under One-Period Forward Guidance



Notes: Graph of the function, $f_I(\pi_I)$, discussed after equation (35). The two crossings with the zero line correspond to equilibria in period I , the date when r switches from $r = r_\ell$ to $r = r_{ss}$. Monetary policy in period I corresponds to one-period forward guidance—that is, the interest rate is held at zero in period I and then reverts to R_{ss} . The red star indicates the level of inflation in period I in the absence of forward guidance.

We now compute the equilibrium allocations in the periods before I_{-1} conditional on the continuation equilibrium starting in period I . The period I_{-1} equilibrium conditions are the appropriate analog of equations (27) through (29):

$$1 = \frac{1}{1 + r_\ell} \left[p \frac{C_\ell}{\pi_\ell C_\ell} + (1 - p) \frac{C_\ell}{\pi_I C_I} \right] \quad (36)$$

²²From equation (33) we see that C_I does not vary with π_I . It follows that f_I is quadratic function of π_I , so that the two solutions displayed in Figure 7 are the only zeros of f_I .

$$\begin{aligned}
& (\pi_\ell - 1) \pi_\ell (C_\ell + G_\ell) - \frac{\varepsilon - 1}{\phi} (\chi N_\ell C_\ell - 1) N_\ell \\
& - \frac{1}{1 + r_\ell} \left[p (\pi_\ell - 1) \pi_\ell (C_\ell + G^\ell) + (1 - p) (\pi_I - 1) \pi_I \frac{C_\ell}{C_I} (C_I + G_{ss}) \right] = 0
\end{aligned} \tag{37}$$

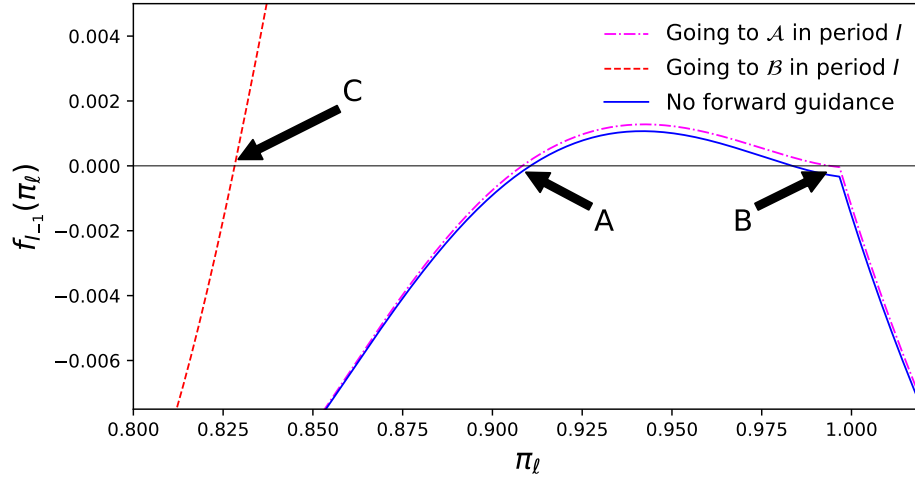
$$N_\ell = (C_\ell + G^\ell) \left(1 + \frac{\phi}{2} (\pi_\ell - 1)^2 \right) \tag{38}$$

Here, we impose the condition that $R_\ell = 1$. In effect, we assume that the ZLB is binding in periods I_{-1} , and the Taylor rule holds. In all of the examples that we have studied, this assumption is satisfied.

We now compute the equilibrium allocations in the periods before I , which we denote by I_{-1} . Given C_I and π_I , equations (36) through (38) define a mapping from π_ℓ to C_ℓ and N_ℓ . Now, we can express the left-hand side of equation (37) as a function of π_ℓ . We denote this function by $f_{I_{-1}}(\pi_\ell; \pi_I, C_I)$. There are two functions, $f_{I_{-1}}$, conditional on the π_I, C_I associated with the period I continuation equilibria, $\mathcal{A}$ and $\mathcal{B}$.

Figure 8 displays both $f_{I_{-1}}$ functions for a range of values of π_ℓ ; see the dotted and dot-dashed lines. We chose the range of π_ℓ so that the graph only displays zeros of $f_{I_{-1}}$ that correspond to equilibria. We find two equilibria corresponding to the $f_{I_{-1}}$ associated with $\mathcal{A}$ (see D and E in Figure 8) and one associated with $\mathcal{B}$ (see C in Figure 8). So there are three REEs with forward guidance. The two REEs without forward guidance can be seen in the solid line in Figure 8 (this curve is taken from Figure 1).

Figure 8: REE Equilibria at the ZLB with and without Forward Guidance



Notes: The solid line reproduces the solid line in Figure 1 and corresponds to the case of no forward guidance. The dashed and dot-dashed lines correspond to the case of forward guidance. The dashed line corresponds to the case in which the economy goes to point $\mathcal{B}$ in the period of the switch in r to r_{ss} (that is, period I). It crosses the zero line more than once, but the other crossing involves very high inflation and is not an equilibrium because the present value of intermediate goods monopolists is negative. The dot-dashed line corresponds to the case in which the economy goes to point $\mathcal{A}$ in period I (see Figure 7).

8.2.2 Learning Equilibria

In the period of forward guidance, $r = r_{ss}$, $R = 1$. In all periods when $r = r_\ell$ (that is, I_{-1}), people understand that the economy reverts to an REE when $r = r_{ss}$. As discussed, there are two REEs starting in period I , the first date when $r = r_{ss}$ (see points $\mathcal{A}$ and $\mathcal{B}$ in Figure 21).

We are interested in three questions. First, do any of the learning equilibria converge to a particular REE in I_{-1} ? Second, if any do converge, how quickly do they do so? Third, are the effects of forward guidance different under learning and rational expectations?

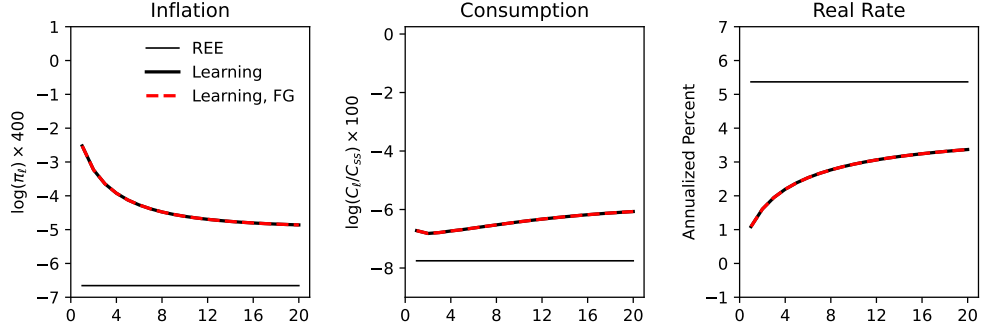
Consider the first question. Two of the three REEs in I_{-1} are not learnable. These are the equilibria associated with points $\mathcal{A}$ and $\mathcal{C}$ in Figure 8. In contrast, the equilibrium represented by $\mathcal{B}$ is learnable. Thus, learnability selects a unique REE.

We now consider a learning equilibrium using the same initial values for Θ as in Section 7.1. Interestingly, forward guidance has no measurable effect on the learning equilibrium. The two learning equilibria are indistinguishable in Figure 9. It follows that the learning equilibrium converges slowly and that there is no forward guidance puzzle under learning.²³ Promises about interest rates in the future do not have implau-

²³For a discussion of the forward guidance puzzle, see Del Negro et al. (2023).

sibly large effects on current economic outcomes. Indeed, under internalized learning, these effects are virtually zero.²⁴

Figure 9: Forward Guidance Under Learning and REE



The power of forward guidance under REE reflects its strong effect on expected inflation. Under learning, the effects of forward guidance have very little influence on expected inflation expectations, because expectations are backward-looking.

9 The Analog of b in the NK Model

In analyzing our nonlinear model, we used the ‘learning principal’ that emerged from the Bray and Savin model: The larger the parameter, b , that controls how self-fulfilling expectations are, the longer it takes to converge. In this section, we accomplish two tasks. First, we demonstrate that the analog of b in the linearized solution of our NK model is the largest real part of the eigenvalues of the matrix that maps beliefs about $x = \begin{bmatrix} C & \pi \end{bmatrix}$ into the realized values of x . Specifically, we develop the analog of Proposition 1 for the NK model, which characterizes the asymptotic rate of convergence of the learning equilibrium as a function of b . Second, we argue that the asymptotic rate of convergence is a good guide to the small t rate of convergence.

We base our analysis below on linearized versions of the policy functions defined in equations (15) and (22). Here, we find it convenient to use time notation rather than recursive notation. The details of our linearization appear in Appendix D. Recall that the household problem can be reduced to finding an optimal decision rule for bond

²⁴Under anticipated utility (not displayed) forward guidance has a slight effect, but not large enough to be economically meaningful.

holdings, $b'(b_h, \Theta, x)$, denoted here by $b_{h,t}$. Log-linearizing this decision rule, we obtain

$$\hat{b}_{h,t} = \gamma_{b,b} \hat{b}_{h,t-1} + \gamma_{b,\pi} \hat{\pi}_t + \gamma_{b,C} \hat{C}_t + \gamma_{b,\mu_\pi} \hat{m}_{\pi,t} + \gamma_{b,\mu_C} \hat{m}_{C,t}. \quad (39)$$

With one exception, the hat notation, $\hat{q}_t$, denotes $\log(q_t/q)$ where q denotes the REE value of q_t about which the linearization is done. The exception is household bond holdings, $b_{h,t}$, in which case $\hat{b}_{h,t}$ denotes $b_{h,t} - b_h$. Also, $\hat{\mu}_t = [\hat{m}_{\pi,t}, \hat{m}_{C,t}]$ represents the log deviation of people's time t posterior of $\mathbb{E}_t x_{t+1}$ and the REE value of $E x_{t+1}$ conditional on $r_{t+1} = r^\ell$.²⁵ We use the posterior means ($\mathbb{E}_t x_{t+1}$) rather than the prior means ($\mathbb{E}_{t-1} x_{t+1}$) because, with Θ and x , households and firms can compute Θ' . Variance of beliefs do not appear because of the certainty equivalence implied by the linearization. The parameters in equation (39) are functions of model parameters and the point about which the linearization is done. These points correspond to different REEs when $r = r^\ell$. Similarly, the linearized price decision rule, $p'_f(p_f, \Theta, x)$, of the firm (denoted by $\hat{p}_{f,t}$) is

$$\hat{p}_{f,t} = \gamma_{p,p} \hat{p}_{f,t-1} + \gamma_{p,\pi} \hat{\pi}_t + \gamma_{p,C} \hat{C}_t + \gamma_{p,\mu_\pi} \hat{m}_{\pi,t} + \gamma_{p,\mu_C} \hat{m}_{C,t}. \quad (40)$$

The time t realized value of $\hat{x}_t$ enters the decision rules, equations (39) and (40), by two channels. The first channel reflects that people use $\hat{x}_t$ to determine the period t values of the exogenous variables in their period t budget constraint. The second channel reflects that $\hat{\mu}_t$ depends on $\hat{x}_t$, $\hat{\mu}_{t-1}$, and the gain in the Bayesian updating equation.

In each period we compute a linearized period equilibrium (see Definition 2), so that (i) $\hat{b}_{h,t-1} = \hat{p}_{f,t-1} = 0$ and (ii) $\hat{x}_t$ is determined by the requirements, $\hat{b}_{h,t} = 0$ and $\hat{p}_{f,t} = 0$:

$$\begin{aligned} 0 &= \gamma_{b,\pi} \hat{\pi}_t + \gamma_{b,C} \hat{C}_t + \gamma_{b,\mu_\pi} \hat{m}_{\pi,t} + \gamma_{b,\mu_C} \hat{m}_{C,t} \\ 0 &= \gamma_{p,\pi} \hat{\pi}_t + \gamma_{p,C} \hat{C}_t + \gamma_{p,\mu_\pi} \hat{m}_{\pi,t} + \gamma_{p,\mu_C} \hat{m}_{C,t}. \end{aligned}$$

²⁵Note that the only exogenous random variable in our model is the natural rate of interest. Given the model structure, in the REE, $E(x_{t+1}|r_{t+1} = r^\ell)$ is equal to the value of x_t while $r_t = r^\ell$. We discuss posterior and prior means because when the natural rate of interest is low households and firms believe that there are additional sources of variation (see equation 25). See Mayer (2021) for a discussion of rates of convergence in univariate learning models with additional stochastic regressors.

Assuming the relevant matrix inverse exists, $\hat{x}_t$ is given by

$$\hat{x}_t = B\hat{\mu}_t, \quad (41)$$

where²⁶

$$B = - \begin{bmatrix} \gamma_{b,\pi} & \gamma_{b,C} \\ \gamma_{p,\pi} & \gamma_{p,C} \end{bmatrix}^{-1} \begin{bmatrix} \gamma_{b,\mu_\pi} & \gamma_{b,\mu_C} \\ \gamma_{p,\mu_\pi} & \gamma_{p,\mu_C} \end{bmatrix}.$$

The law of motion of $\hat{\mu}_t$ is a stacked version of the updating expressions in equation (4). For simplicity, we impose the same gain, $\gamma_t = 1/(\lambda_0 + t)$, on the two equations. Here, λ_0 denotes the initial precision of beliefs about the mean of inflation and consumption. Combining the vector Bayesian updating expression with equation (41) we obtain:

$$\hat{\mu}_t = (I - B_t) \hat{\mu}_{t-1}, \quad (42)$$

where $B_t = \gamma_t [I - B(1 - \gamma_t)(I - \gamma_t B)^{-1}]$ is the analog of b_t in equation (8).²⁷ A difference is that $(1 - \gamma_t)(I - \gamma_t B)^{-1}$ does not appear in b_t , reflecting the timing differences between the two models. These timing differences are negligible for our purpose; if we replaced $(1 - \gamma_t)(I - \gamma_t B)^{-1}$ by I , then the asymptotic convergence result stated below would be unchanged.

The mapping from beliefs about $\hat{x}_t$ —that is, $\hat{\mu}_{t-1}$ —to realized values of $\hat{x}_t$ is obtained by multiplying equation (42) by B and using equation (41):

$$\hat{x}_t = B(I - B_t) \hat{\mu}_{t-1}. \quad (43)$$

The period equilibrium of the linearized model corresponds to the values of $\hat{x}_t, \hat{\mu}_t, \gamma_t$ which solve equations (42) and (43).

For t large enough, equation (43) is approximately $\hat{x}_t = B\hat{\mu}_{t-1}$. In a slight abuse of notation, we let b denote the largest real part of the eigenvalues of B . The central result of this subsection is that b characterizes the asymptotic speed at which the learning equilibrium converges to the stable REE. Therefore, it plays the same role as

²⁶In the examples that we have considered, we have not encountered an exception to the invertibility assumption.

²⁷According to the Bayesian updating equations, $\hat{\mu}_t = \hat{\mu}_{t-1} + \gamma_t(\hat{x}_t - \hat{\mu}_{t-1})$. Substituting out for $\hat{x}_t$ using equation (41), we obtain $\hat{\mu}_t = \hat{\mu}_{t-1} + \gamma_t(B\hat{\mu}_t - \hat{\mu}_{t-1})$. Equation (42) follows after rearranging.

the parameter, b , in the Bray and Savin (1986) model.

Our analysis of the speed of convergence of $\hat{\mu}_t$ holds for any finite-dimensional $\hat{\mu}_t$. We extend Definition 1 to the vector case of $\hat{x}_t$ as follows:²⁸

Definition 4. The vector series, $\hat{x}_t$, and the scalar series, $a_t > 0$, converge to zero at the same rate if (i) for all j either (a) $\hat{x}_{j,t} \simeq a_t$ or (b) $\lim_{t \rightarrow \infty} |\hat{x}_{j,t}|/a_t = 0$, and (ii) for at least one j , condition (a) holds.

Our definition requires that all elements of $\hat{x}_t$ converge to zero at least as fast as a_t and that at least one element converges to zero at the same rate as a_t .

In what follows, it is useful to denote the eigenvalue-eigenvector decomposition of B as $B = Q\Lambda Q^{-1}$, where Λ is a diagonal matrix with the eigenvalues of B in the diagonal elements.²⁹ In the NK model, one cannot rule out the possibility that the eigenvalues, Λ_j , of B are complex. Let $\Lambda_j = \Lambda_{r,j} + i\Lambda_{c,j}$, where $\Lambda_{r,j}$ and $\Lambda_{c,j}$ denote the real and complex parts of Λ_j , respectively. Let $r_{j,t}$ denote the modulus of the j^{th} eigenvalue of $I - B_t$, and let j^* denote a value of j for which $\Lambda_{r,j}$ attains its maximal value, b . Also, define $\tilde{\mu}_t = Q^{-1}\hat{\mu}_t$, where $\hat{\mu}_0$ is given. The following proposition establishes the rate of convergence of $\hat{\mu}_t$.

Proposition 2. *Suppose that (i) B has an eigenvalue-eigenvector decomposition with $\Lambda_{r,j} < 1$ for all j , (ii) $\tilde{\mu}_{j^*,0} \neq 0$, and (iii) $r_{j^*,t} \neq 0$ for each t . Then, $\hat{\mu}_t \simeq t^{b-1}$.*

See Appendix A for a proof.

Some comments about Proposition 2 are in order. First, violations of (ii) or (iii) are isolated special cases. Condition (ii) rules out the case in which the initial priors are orthogonal to the left eigenvector associated with the eigenvalue, Λ_{j^*} . In that case, Λ_{j^*} plays no role in the system's dynamics. Condition (iii) is analogous to the requirement in Proposition 1 that $\gamma_t(1 - b) \neq 1$. Second, our proposition is consistent with the result of Evans and Honkapohja (2001, Theorem 1) that $\lim_{t \rightarrow \infty} \hat{\mu}_t = \mathbf{0}$. The novelty of Proposition 2 is that it establishes the rate at which $\hat{\mu}_t$ converges to zero. Third, as in our analysis of the Bray and Savin (1986) model, the rate of convergence of learning in the NK model is decreasing in b . Fourth, the fact that b plays a similar role in the Bray and Savin (1986) and NK models can be seen by noting that pre-multiplication of equation (42) by Q^{-1} diagonalizes the system into a set of first-order

²⁸See Definition 1 for ' $\simeq$ '.

²⁹A sufficient condition for this decomposition to exist is that the eigenvalues of B are distinct. This condition is satisfied in all the examples that we consider.

scalar difference equations in $\tilde{\mu}_{j,t}$ that are independent across j . Each of these equations resembles equation (6) in Bray and Savin (1986). So the representation of $\tilde{\mu}_{j,t}$ has the form given in equations (7) and (9) (though $\tilde{\mu}_{j,t}$ is potentially complex valued), and its behavior is determined by the j^{th} eigenvalue of B . Since $\hat{\mu}_t = Q\tilde{\mu}_t$ and the columns of Q are linearly independent, it follows that the largest real part of the eigenvalues of B determines the rate of convergence of at least one element of $\hat{\mu}_t$. Fifth, when the eigenvalues of B are complex, $\hat{\mu}_t$ can exhibit sinusoidal fluctuations. That is why our definition of the rate of convergence (Definition 1) needs to accommodate the possibility that the $\hat{\mu}_t/t^{b-1}$ oscillates in a bounded set.³⁰

Table 2 displays the eigenvalues of B corresponding to the Good-ZLB and Bad-ZLB equilibria for the benchmark parameter values. The maximal eigenvalue (‘Eigenvalue 1’), b , associated with the Good-ZLB and Bad-ZLB equilibria are 0.92 and 1.26, respectively. Consistent with the claim in Section 6, the Bad-ZLB equilibrium is not locally learnable because $b > 1$. The Good-ZLB equilibrium is locally learnable because, in that case, $b < 1$.

Table 2 displays the eigenvalues of B corresponding to the Good- and Bad-ZLB equilibria for the benchmark parameter values. The maximal eigenvalue (‘Eigenvalue 1’), b , associated with the Good-ZLB and Bad-ZLB equilibria are 0.92 and 1.26, respectively. Consistent with the claim in Section 6, the Bad-ZLB equilibrium is not locally learnable because $b > 1$. The Good-ZLB equilibrium is locally learnable because, in that case, $b < 1$.

Asymptotic convergence to the Good-ZLB REE is slow because b is large. According to Proposition 2 $|\hat{\mu}_{j^*,t}|$ is approximately κt^{b-1} for some $\kappa \neq 0$ and t sufficiently large. The amount of time it takes to close two-thirds of a gap, $\hat{\mu}_t$, for t sufficiently large, is given by T_t in equation (11). Table 2 reports values of T_1 for different variants of the model. In the benchmark model, when $b = 0.92$, then $T_1 = 920,482$. This large value of T_1 is qualitatively consistent with the basic prediction of the nonlinear solution to the model—namely, that the rate of convergence is quite slow (see Figure 4a). Similarly, the small value of T_1 reported in the table for the case in which the ZLB is not binding and $\alpha = 1.5$ is qualitatively consistent with the finding for the nonlinear solution to the model (see Figure 5). In this sense, the asymptotic result in Proposition 2 is a useful guide about the rate of convergence, even for small t .

A different way to assess the usefulness of the asymptotics is to calculate the actual

³⁰A discussion about the possibility of oscillations follows the proof of Proposition 2 in Appendix A.

amount of time, T , required to close two-thirds of the initial gap between priors and steady state according to the linearized solution to the model.³¹ To this end, we simulate the linearized solution to the model when the ZLB is binding and when we ignore the ZLB. In the latter case, we consider $\alpha = 1.5$ and 3. The results are reported in Table 4a. We find that, for the benchmark model, when the ZLB is binding, $T = 944,710$. In sharp contrast, when the ZLB is not binding and $\alpha = 1.5$ and 3, we find that $T = 3$ and 1 periods, respectively. These results about the importance of the ZLB and the value of α in determining the speed of convergence are qualitatively the same as our results using T_1 . Thus, the rule of thumb, equation (11), suggested by Proposition 2 is informative about actual rates of convergence in the linearized solution to the model.

Table 2: Eigenvalues of B

	Eigenvalue 1	Eigenvalue 2	T_1	T
Good ZLB	0.92	-0.48	920,482	944,710
Bad ZLB	1.26	-1.21	NA	NA
No ZLB, $\alpha = 1.5$	0.054+0.44i	0.054-0.44i	2	3
No ZLB, $\alpha = 3$	-0.135+0.84i	-0.135-0.84i	2	1

Note: The matrix, B , is defined in equation (41). The scalar, b , discussed in the text is the largest real part of the eigenvalues of B . The reported values of T are based on simulations of the linearized solution to the model. For the definitions of T and T_1 see the text.

10 Conclusion

In this paper, we consider the speed with which people learn about their environment after an unusual event. We do so in a non-linear NK model with internally rational households and firms that are learning about how the economy will evolve after the event. To characterize the speed of convergence of people's beliefs, we analytically extend results in the literature to encompass circumstances when learning is very slow. We argue that the slow convergence result arises naturally in the NK model when the ZLB is binding. Under these circumstances, analyses of fiscal and monetary policies under rational expectations can be very misleading. Since inflation declined by a modest amount during the Great Recession, learning moves the model toward the data relative to rational expectations. In this sense, learning provides a possible resolution

³¹The initial gap in $\log x_i$, $i = 1, 2$, corresponds to the log-deviation of x_i in the initial steady state and the REE while $r = r_\ell$.

to the ‘missing deflation puzzle’ (see Del Negro et al. (2023)). It would be interesting to pursue this possibility in an empirically plausible version of the NK model of the sort considered by Christiano et al. (2016) or Del Negro et al. (2023).

Finally, we note that there are other circumstances in which slow learning could arise. For example, plausible parameterizations of Cagan (1956)’s model of money demand under hyperinflation map into high b economies. Results in Marcet and Sargent (1995, Table 6.3) imply that estimates of the elasticity of money demand in hyperinflations (see, for example, Christiano (1987) and Taylor (1991)) map into high values of b . More generally, the learning principle suggests that any model with strong strategic complementarities could exhibit slow convergence to rational expectations under learning.

References

- Adam, Klaus, Albert Marcet, and Johannes Beutel**, “Stock price booms and expected capital gains,” *American Economic Review*, 2017, *107* (8), 2352–2408.
- **and —**, “Internal rationality, imperfect market knowledge and asset prices,” *Journal of Economic Theory*, 2011, *146* (3), 1224–1252.
- **and Sebastian Merkel**, “Stock price cycles and business cycles,” *European Central Bank Working Paper no. 2316*, September 2019.
- Angeletos, George-Marios and Chen Lian**, “Dampening general equilibrium: From micro to macro,” Technical Report, National Bureau of Economic Research 2017.
- **and —**, “Forward Guidance without Common Knowledge,” *American Economic Review*, September 2018, *108* (9), 2477–2512.
- Arifovic, Jasmina, Stephanie Schmitt-Grohé, and Martin Uribe**, “Learning to live in a liquidity trap,” *Journal of Economic Dynamics and Control*, 2018, *89*, 120–136.
- Benhabib, Jess, Stephanie Schmitt-Grohé, and Martin Uribe**, “Monetary policy and multiple equilibria,” *American Economic Review*, 2001, *74* (2), 165–170.
- Benveniste, Albert, Michel Métivier, and Pierre Priouret**, *Adaptive algorithms and stochastic approximations*, Vol. 22, Springer Science & Business Media, 1990.
- Bray, Margaret M and Nathan E Savin**, “Rational expectations equilibria, learning, and model specification,” *Econometrica*, 1986, pp. 1129–1160.
- Cagan, Phillip**, “The monetary dynamics of hyperinflation,” *Studies in the Quantity Theory of Money*, 1956.
- Chien, YiLi, In-Koo Cho, and B Ravikumar**, “Convergence to Rational Expectations in Learning Models: A Note of Caution,” *Available at SSRN 3885395*, 2021.
- Christiano, Lawrence J**, “Cagan’s model of hyperinflation under rational expectations,” *International Economic Review*, 1987, pp. 33–49.
- Christiano, Lawrence J., Martin S. Eichenbaum, and Mathias Trabandt**, “Unemployment and Business Cycles,” *Econometrica*, 2016, *84* (4), 1523–1569.
- , — , **and Sergio Rebelo**, “When is the Government Spending Multiplier Large?,” *Journal of Political Economy*, 2011, (February).
- Christopeit, Norbert and Michael Massmann**, “Estimating structural parameters in regression models with adaptive learning,” *Econometric Theory*, 2018, *34* (1), 68–111.

- Cogley, Timothy and Thomas J Sargent**, “Anticipated utility and rational expectations as approximations of Bayesian decision making,” *International Economic Review*, 2008, *49* (1), 185–221.
- Eggertsson, Gauti and Michael Woodford**, “The Zero Bound on Interest Rates and Optimal Monetary Policy,” *Brookings Papers on Economic Activity*, 2003, (1).
- Eggertsson, Gauti B and Michael Woodford**, “Policy options in a liquidity trap,” *American Economic Review*, 2004, *94* (2), 76–79.
- Erceg, Christopher J. and Andrew T. Levin**, “Imperfect credibility and inflation persistence,” *Journal of Monetary Economics*, 2003, *50* (4), 915–944.
- Eusepi, Stefano, Christopher Gibbs, and Bruce Preston**, “Monetary policy trade-offs at the zero lower bound,” *Available at SSRN 4069078*, 2022.
- Evans, George W**, “Theories of learning and economic policy,” *Revue d’économie politique*, 2021, *132* (3-4), 583–608.
- **and Seppo Honkapohja**, “Convergence for difference equations with vanishing time-dependence, with applications to adaptive learning,” *Economic Theory*, 2000, *15*, 717–725.
- Evans, George W. and Seppo Honkapohja**, *Learning and Expectations in Macroeconomics*, Princeton University Press, 2001.
- Farhi, E, M Petri, and I Werning**, “The fiscal multiplier puzzle: Liquidity traps, bounded rationality, and incomplete markets,” Technical Report, mimeo 2020.
- Farhi, Emmanuel and Iván Werning**, “Monetary policy, bounded rationality, and incomplete markets,” *American Economic Review*, 2019, *109* (11), 3887–3928.
- Farmer, Leland, Emi Nakamura, and Jón Steinsson**, “Learning about the long run,” Technical Report, National Bureau of Economic Research 2021.
- Ferrero, Giuseppe**, “Monetary policy, learning and the speed of convergence,” *Journal of Economic Dynamics and Control*, 2007, *31* (9), 3006–3041.
- García-Schmidt, Mariana and Michael Woodford**, “Are low interest rates deflationary? A paradox of perfect-foresight analysis,” *American Economic Review*, 2019, *109* (1), 86–120.
- Gust, Christopher, Edward Herbst, and David López-Salido**, “Forward guidance with bayesian learning and estimation,” 2018.
- Hamilton, James D**, *Time series analysis*, Princeton university press, 2020.

- Heemeijer, Peter, Cars Hommes, Joep Sonnemans, and Jan Tuinstra**, “Price stability and volatility in markets with positive and negative expectations feedback: An experimental investigation,” *Journal of Economic dynamics and control*, 2009, *33* (5), 1052–1072.
- Hommes, Cars**, “The heterogeneous expectations hypothesis: Some evidence from the lab,” *Journal of Economic dynamics and control*, 2011, *35* (1), 1–24.
- Iovino, Luigi and Dmitriy Sergeyev**, “Central bank balance sheet policies without rational expectations,” *Review of Economic Studies*, 2023, *90* (6), 3119–3152.
- Kaplan, Greg and Giovanni L Violante**, “Microeconomic heterogeneity and macroeconomic shocks,” *Journal of Economic Perspectives*, 2018, *32* (3), 167–94.
- Kreps, David M**, “Anticipated utility and dynamic choice,” *Econometric Society Monographs*, 1998, *29*, 242–274.
- Ljung, Lennart**, “Analysis of recursive stochastic algorithms,” *IEEE transactions on automatic control*, 1977, *22* (4), 551–575.
- Lucas, Robert E**, “Some international evidence on output-inflation tradeoffs,” *The American economic review*, 1973, pp. 326–334.
- Lucas, Robert E.**, “Asset Prices in an Exchange Economy,” *Econometrica*, 1978, *46* (6), 1429–1445.
- Marcet, Albert and Thomas J Sargent**, “Convergence of least-squares learning in environments with hidden state variables and private information,” *Journal of Political Economy*, 1989, *97* (6), 1306–1322.
- **and** —, “Convergence of least squares learning mechanisms in self-referential linear stochastic models,” *Journal of Economic theory*, 1989, *48* (2), 337–368.
- **and** —, “Convergence of least squares learning mechanisms in self-referential linear stochastic models,” *Journal of Economic theory*, 1989, *48* (2), 337–368.
- **and Thomas Sargent**, “Speed of Convergence of Recursive Least Squares: Learning with Autoregressive Moving-Average Perceptions,” in Alan Kirman and Mark Salmon, eds., *Learning and rationality in economics*, Oxford University Press, UK, 1995, chapter 6, pp. 179–215.
- Muth, John F**, “Rational expectations and the theory of price movements,” *Econometrica*, 1961, pp. 315–335.
- Negro, Marco Del, Marc P Giannoni, and Christina Patterson**, “The forward guidance puzzle,” *Journal of Political Economy Macroeconomics*, 2023, *1* (1), 43–79.

- Preston, Bruce**, “Adaptive learning in infinite horizon decision problems,” *Unpublished, Columbia University*, 2005.
- Rotemberg, Julio J**, “Sticky prices in the United States,” *Journal of political economy*, 1982, *90* (6), 1187–1211.
- Taylor, Mark P**, “The hyperinflation model of money demand revisited,” *Journal of money, Credit and Banking*, 1991, *23* (3), 327–351.
- Vimercati, Riccardo Bianchi, Martin S Eichenbaum, and Joao Guerreiro**, “Fiscal policy at the zero lower bound without rational expectations,” Technical Report, National Bureau of Economic Research 2021.
- Vives, Xavier**, “How fast do rational agents learn?,” *The Review of Economic Studies*, 1993, *60* (2), 329–347.
- Werning, Ivan**, “Managing a Liquidity Trap: Monetary and Fiscal Policy, manuscript,” 2012.
- Woodford, Michael**, “Self-fulfilling expectations and fluctuations in aggregate demand,” 1990.
- , “Methods of Policy Accommodation at the Interest-rate Lower Bound,” *Proceedings-Economic Policy Symposium-Jackson Hole, Federal Reserve Bank of Kansas City*, 2012, pp. 185–288.
- **and Yixi Xie**, “Policy options at the zero lower bound when foresight is limited,” in “AEA Papers and Proceedings,” Vol. 109 American Economic Association 2014 Broadway, Suite 305, Nashville, TN 37203 2019, pp. 433–437.
- **and —**, “Fiscal and monetary stabilization policy at the zero lower bound: Consequences of limited foresight,” *Journal of Monetary Economics*, 2022, *125*, 18–35.

A Proofs of Lemmas and Propositions

A.1 Proof of Lemma 1

Lemma 1 follows from the following Proposition.

Proposition 3. *Suppose $b_t \neq 1$ for all t , then μ_t can be written as*

$$\mu_t = \frac{a}{1-b} + \sum_{j=1}^t \left\{ \frac{z_t}{z_j} b_j \frac{\varepsilon_j}{1-b} \right\} + z_t \left(\mu_0 - \frac{a}{1-b} \right)$$

and has mean

$$\frac{a}{1-b} + z_t \left(\mu_0 - \frac{a}{1-b} \right)$$

and variance

$$\sum_{j=1}^t \left\{ \left(\frac{z_t}{z_j} \right)^2 b_j^2 \frac{\sigma_\varepsilon^2}{(1-b)^2} \right\}.$$

Proof. Note that

$$\begin{aligned} \left(\mu_t - \frac{a}{1-b} \right) &= b_t \frac{\varepsilon_t}{1-b} + (1-b_t) \left(\mu_{t-1} - \frac{a}{1-b} \right) \\ &= b_t \frac{\varepsilon_t}{1-b} + (1-b_t) b_{t-1} \frac{\varepsilon_{t-1}}{1-b} + (1-b_t) (1-b_{t-1}) \left(\mu_{t-2} - \frac{a}{1-b} \right) \\ &= b_t \frac{\varepsilon_t}{1-b} + (1-b_t) b_{t-1} \frac{\varepsilon_{t-1}}{1-b} + (1-b_t) (1-b_{t-1}) b_{t-2} \frac{\varepsilon_{t-2}}{1-b} \\ &\quad + (1-b_t) (1-b_{t-1}) (1-b_{t-2}) \left(\mu_{t-3} - \frac{a}{1-b} \right) \\ &= b_t \frac{\varepsilon_t}{1-b} + \sum_{j=1}^{t-1} \left\{ \left[\prod_{k=1}^j (1-b_{t-k+1}) \right] b_{t-j} \frac{\varepsilon_{t-j}}{1-b} \right\} + z_t \left(\mu_0 - \frac{a}{1-b} \right) \\ &= \sum_{j=1}^t \left\{ \frac{z_t}{z_j} b_j \frac{\varepsilon_j}{1-b} \right\} + z_t \left(\mu_0 - \frac{a}{1-b} \right). \end{aligned}$$

The results of the proposition follow immediately from the properties of ε_t . □

A.2 Proof of Proposition 1

We first state a number of lemmas that will be useful in the proof.

Lemma 2. For any $b < 1$ and any $0 \leq \lambda_0$, there exists a t^* so that $0 < b_t < 1$ for all $t \geq t^*$.

Proof. Let $t^* = \max \{1, \lceil 2 - b - \lambda_0 \rceil\}$ where $\lceil x \rceil$ is the smallest integer larger than x . The result follows immediately. $\square$

Lemma 3. Define $H_{\lambda_0, T} = \sum_{t=1}^T \frac{1}{\lambda_0 + t}$. Suppose $0 \leq \lambda_0 < \infty$, then

$$\lim_{T \rightarrow \infty} \{H_{\lambda_0, T} - \log(\lambda_0 + T)\} = c_{\lambda_0},$$

where c_{λ_0} is a finite constant.

Proof. Note that $H_{\lambda_0, T} > 0$ and

$$H_{\lambda_0, T} \leq \frac{1}{\lambda_0 + 1} + \int_{\lambda_0 + 1}^{\lambda_0 + T} \frac{1}{t} dt = \frac{1}{\lambda_0 + 1} + \log(\lambda_0 + T) - \log(\lambda_0 + 1).$$

Define the sequence $y_{\lambda_0, T} = \log(\lambda_0 + 1) + H_{\lambda_0, T} - \log(\lambda_0 + T)$. The above inequality and the convexity of t^{-1} imply $0 < y_{\lambda_0, T} \leq \frac{1}{\lambda_0 + 1}$ for all $T \geq 1$. Also,

$$y_{\lambda_0, T+1} - y_{\lambda_0, T} = \frac{1}{\lambda_0 + T + 1} + \log(\lambda_0 + T) - \log(\lambda_0 + T + 1) < 0$$

Thus, $y_{\lambda_0, T}$ is a monotone, decreasing series. The result of the lemma follows by the monotone convergence theorem. $\square$

Lemma 4. Define $H_{\lambda_0, T} = \sum_{t=1}^T \frac{1}{\lambda_0 + t}$. Suppose $0 \leq \lambda_0 < \infty$, then there exist positive, finite constants $\underline{c}_{\lambda_0}$ and $\bar{c}_{\lambda_0}$, and a finite constant c_{λ_0} so that

$$\frac{\underline{c}_{\lambda_0}}{T} < \log(T + \lambda_0) + c_{\lambda_0} - H_{\lambda_0, T} < \frac{\bar{c}_{\lambda_0}}{T}.$$

Proof. Using c_{λ_0} from Lemma 3, noting that t^{-1} is convex, and following the geometric logic in Young (1991) “Euler’s Constant” The Mathematical Gazette, Vol. 75, No. 472 (Jun., 1991) we have

$$\frac{1}{2(\lambda_0 + T)} < \log(\lambda_0 + T) + c_{\lambda_0} - H_{\lambda_0, T} < \frac{1}{\lambda_0 + T}.$$

The result of the lemma follows immediately. $\square$

We are now in a position to prove Proposition 1.

Proof. Consider $0 < b < 1$. Note that $0 < b_j < 1$ for all j . Define $y_t = \log \left((\lambda_0 + t)^{1-b} z_t \right)$. Note that,

$$\begin{aligned} y_{t+h} - y_t &= (1-b) \log \left(1 + \frac{h}{\lambda_0 + t} \right) + \sum_{k=1}^h \log (1 - b_{t+k}) \\ &= (1-b) \left(\log \left(1 + \frac{h}{\lambda_0 + t} \right) - \sum_{k=1}^h \frac{1}{\lambda_0 + t + k} \right) + R_{t,t+h} \end{aligned}$$

where $R_{t,t+h}$ is the remainder term from Taylor's Theorem in the representation of $\log(\cdot)$. By Lemma 3, $\log(\lambda_0 + t) - \sum_{k=1}^t \frac{1}{\lambda_0 + t}$ is a Cauchy sequence. So, for any $\epsilon > 0$ there exists a t_1 so that if $t \geq t_1$ then for any $h \geq 0$

$$\left| \log \left(1 + \frac{h}{\lambda_0 + t} \right) - \sum_{k=1}^h \frac{1}{\lambda_0 + t + k} \right| < \frac{\epsilon}{4 - 4b}.$$

From Taylor's Theorem, $|R_{h,t+h}| \leq \sum_{k=1}^h K \left(\frac{1}{\lambda_0 + t + k} \right)^2$ for some $0 < K < \infty$. Because $\sum_{t=1}^{\infty} t^{-2}$ converges, for any $\epsilon > 0$ there exists a t_2 so that if $t \geq t_2$ then for any $h \geq 0$ $|R_{h,t+h}| < \epsilon/4$. It follows that if $t \geq \max\{t_1, t_2\}$, then $|y_{t+h} - y_t| < \epsilon/2$. As a result, if $t \geq \max\{t_1, t_2\}$ for any $j, k > 0$, $|y_{t+j} - y_{t+k}| < \epsilon$. That is y_t is a Cauchy sequence. The conclusion of the proposition follows for $0 < b < 1$ is then immediate.

Now consider $b \leq 0$. By Lemma 2 there exists a t^* so that for all $t \geq t^*$ we have $0 < b_t < 1$. Define the sequence $y_t^* = (1-b) \log(\lambda_0 + t) + \sum_{j=t^*}^t \log(1 - b_j)$. Using the same argument as in the case when $0 < b < 1$, we have that y_t^* is Cauchy, so it converges to a finite constant. For $t \geq t^*$, we have

$$t^{1-b} z_t = \left(\frac{t}{\lambda + t} \right)^{1-b} \exp(y_t^*) \left[\prod_{j=1}^{t^*-1} (1 - b_j) \right].$$

The conclusion of the proposition for $b \leq 0$ follows by noting that we have ruled out the possibility that $\prod_{j=1}^{t^*-1} (1 - b_j) = 0$ through our assumption that $b_t \neq 1$ for all t . $\square$

A.3 Proof of Proposition (2) and related results

Throughout, we assume that B has an eigenvalue-eigenvector decomposition meaning that $B = Q\Lambda Q^{-1}$ where the columns of Q are linearly independent eigenvectors of B and Λ is a diagonal matrix with eigenvalue $\Lambda_{j,r} + \Lambda_{j,c}i$ in the j th diagonal element. To

prove the proposition, we will need some Lemmas and notation.

Lemma 5. $\Omega_t = B \left(I - \frac{1}{\lambda_0 + t} B \right)^{-1} \left(1 - \frac{1}{\lambda_0 + t} \right)$ has an eigenvalue-eigenvector decomposition so that $\Omega_t = Q \Lambda_t Q^{-1}$ where Q is the same matrix as in the eigenvector-eigenvalue decomposition of B .

Proof. From the definition of Ω_t

$$\begin{aligned} \Omega_t Q &= B \left(I - \frac{1}{\lambda_0 + t} B \right)^{-1} \left(1 - \frac{1}{\lambda_0 + t} \right) Q \\ &= Q \Lambda Q^{-1} \left(Q Q^{-1} - \frac{1}{\nu_0 + t} Q \Lambda Q^{-1} \right)^{-1} \left(1 - \frac{1}{\nu_0 + t} \right) Q \\ &= Q \Lambda Q^{-1} Q \left(I - \frac{1}{\nu_0 + t} \Lambda \right)^{-1} Q^{-1} \left(1 - \frac{1}{\nu_0 + t} \right) Q \\ &= Q \Lambda_t \end{aligned}$$

where $\Lambda_t = \Lambda \left(I - \frac{1}{\nu_0 + t} \Lambda \right)^{-1} \left(1 - \frac{1}{\nu_0 + t} \right)$ is a diagonal matrix. $\square$

Lemma 5 says that the columns of Q are eigenvectors of Ω_t , though Ω_t has different eigenvalues than B . Denote the j th diagonal element of Λ_t by $\Lambda_{j,r,t} + \Lambda_{j,c,t}i$. The following Lemma follows from tedious algebra.

Lemma 6. Suppose $\Lambda_{j,r,t}$, $\Lambda_{j,r}$, $\Lambda_{j,c,t}$, and $\Lambda_{j,c}$ are as defined above and that for all j either $\frac{1 - \Lambda_{r,j}}{\lambda_0 + t} \neq 1$ or $\Lambda_{c,j,t} \neq 0$ for all t . Then $\lim_{t \rightarrow \infty} \Lambda_{j,r,t} = \Lambda_{j,r}$, $\lim_{t \rightarrow \infty} \Lambda_{j,c,t} = \Lambda_{j,c}$ and there exists a finite, positive constant K so that $|\Lambda_{j,r,t} - \Lambda_{j,r}| < K (\lambda_0 + t)^{-1}$ and $|\Lambda_{j,c,t} - \Lambda_{j,c}| < K (\lambda_0 + t)^{-1}$.

It will be useful to define $\tilde{\mu}_t \equiv Q^{-1} \mu_t$, which is given by

$$\tilde{\mu}_t = \left(I - \frac{1}{\lambda_0 + t} (I - \Lambda_t) \right) \tilde{\mu}_{t-1} = \left(\prod_{k=1}^t \left(I - \frac{1}{\lambda_0 + k} (I - \Lambda_k) \right) \right) \tilde{\mu}_0. \quad (44)$$

The modulus of the j th diagonal element of $\left(I - \frac{1}{\lambda_0 + t} (I - \Lambda_t) \right)$ is

$$r_{j,t} = \left(\left[1 - \frac{1}{\lambda_0 + t} (1 - \Lambda_{j,r,t}) \right]^2 + \left(\frac{1}{\lambda_0 + t} \Lambda_{j,c,t} \right)^2 \right)^{1/2} \quad (45)$$

We will use the following Lemma.

Lemma 7. Suppose $r_{j,t} \neq 0$ for all t and $\Lambda_{j,r} < 1$, then

$$t^{1-\Lambda_{j,r}} \prod_{k=1}^t r_{j,k} \rightarrow \kappa_j > 0$$

where $r_{j,k}$ is given by equation (45).

Proof. For a given j , define

$$y_{j,t} \equiv (\lambda_0 + t)^{1-\Lambda_{j,r}} \prod_{k=1}^t r_{j,k}$$

$$\tilde{y}_{j,t} \equiv \log(y_{j,t}).$$

Note that

$$\begin{aligned} \tilde{y}_{j,t+h} - \tilde{y}_{j,t} &= (1 - \Lambda_{j,r}) \log \left(1 + \frac{h}{\lambda_0 + t} \right) \\ &\quad + \frac{1}{2} \sum_{k=1}^h \log \left(1 + 2 \frac{1}{\lambda_0 + t + k} (\Lambda_{j,r,t+k} - 1) \right. \\ &\quad \left. + \left[\frac{1}{\lambda_0 + t + k} (\Lambda_{j,r,t+k} - 1) \right]^2 + \left(\frac{1}{\lambda_0 + t + k} \Lambda_{j,c,t+k} \right)^2 \right) \\ &= (1 - \Lambda_{j,r}) \left(\log \left(1 + \frac{h}{\lambda_0 + t} \right) - \sum_{k=1}^h \frac{1}{\lambda_0 + t + k} \right) + R_{t,t+h} \\ &\quad + \sum_{k=1}^h \left(\frac{1}{\lambda_0 + t + k} (\Lambda_{j,r,t+k} - \Lambda_{j,r}) \right. \\ &\quad \left. + \frac{1}{2} \left[\frac{1}{\lambda_0 + t + k} (\Lambda_{j,r,t+k} - 1) \right]^2 + \frac{1}{2} \left(\frac{1}{\lambda_0 + t + k} \Lambda_{j,c,t+k} \right)^2 \right) \end{aligned}$$

where $R_{t,t+h}$ is the remainder term from Taylor's theorem. For any ϵ , there exists a t_1 (which does not depend on h) so that if $t \geq t_1$ then for any $k \geq 0$ by Lemmas 3 and 6

we have

$$\begin{aligned}
\left| \log \left(1 + \frac{k}{\nu_0 + t} \right) - \sum_{n=1}^k \frac{1}{\nu_0 + t + n} \right| &< \frac{\epsilon}{6 |1 - \Lambda_{j,r}|} \\
|\Lambda_{j,r,t+k} - \Lambda_{j,r}| &< \frac{K}{\lambda_0 + t + k} \\
(\Lambda_{j,r,t+k} - 1)^2 &< (1 - \Lambda_{j,r})^2 + \epsilon \\
\Lambda_{j,c,t+k}^2 &< \Lambda_{j,c}^2 + \epsilon.
\end{aligned}$$

for some finite, positive K . Then

$$|\tilde{y}_{t+h} - \tilde{y}_t| \leq \frac{\epsilon}{6} + \sum_{k=1}^h \left(\frac{1}{\nu_0 + t + k} \right)^2 \left(K + \frac{1}{2} ((\Lambda_{j,r} - 1)^2 + \epsilon) + \frac{1}{2} (\Lambda_{j,c}^2 + \epsilon) \right) + |R_{t,t+h}|$$

Because $\sum_{t=1}^{\infty} t^{-2}$ converges, for any ϵ there exists a t_2 (which does not depend on h) so that if $t \geq t_2 \geq t_1$ then

$$\sum_{k=1}^{\infty} \left(\frac{1}{\nu_0 + t + k} \right)^2 \left(K + \frac{1}{2} ((\Lambda_{j,r} - 1)^2 + \epsilon) + \frac{1}{2} (\Lambda_{j,c}^2 + \epsilon) \right) < \frac{\epsilon}{6}.$$

Now, consider

$$\begin{aligned}
|R_{t,t+h}| &= \sum_{k=0}^h \frac{1}{(1 + x_k)^2} \frac{1}{4} \left[2 \frac{1}{\lambda_0 + t + k} (\Lambda_{j,r,t+k} - 1) \right. \\
&\quad \left. + \left(\frac{1}{\lambda_0 + t + k} (\Lambda_{j,r,t+k} - 1) \right)^2 + \left(\frac{1}{\lambda_0 + t + k} \Lambda_{j,c,t+k} \right)^2 \right]^2
\end{aligned}$$

for

$$x_k \in \left[0, 2 \frac{1}{\lambda_0 + t + k} (\Lambda_{j,r,t+k} - 1) + \left(\frac{1}{\lambda_0 + t + k} (\Lambda_{j,r,t+k} - 1) \right)^2 + \left(\frac{1}{\lambda_0 + t + k} \Lambda_{j,c,t+k} \right)^2 \right].$$

For any h , if $t \geq t_2$

$$|R_{t,t+h}| < \sum_{k=0}^h \left(\frac{1}{\lambda_0 + t + k} \right)^2 \frac{1}{\left(1 - 2 \frac{1}{\lambda_0 + t_2} (\Lambda_{j,r} - 1 - \epsilon) \right)^2} \frac{1}{4} \\ \times \left[2K + 2(1 - \Lambda_{j,r} + \epsilon) + [(\Lambda_{j,r} - 1)^2 + \epsilon] + (\Lambda_{j,c}^2 + \epsilon) \right]^2.$$

Because $\sum_t t^{-2}$ converges to a finite constant, for any ϵ , there exists a $t_3 \geq t_2$ (which does not depend on h) so that if $t \geq t_3$ then

$$|R_{t,t+h}| < \frac{\epsilon}{6}.$$

Combining results, for any $\epsilon > 0$, there exists a t_3 so that if $t \geq t_3$, then for any h

$$|\tilde{y}_{t+h} - \tilde{y}_t| < \frac{\epsilon}{2}.$$

Then for any $n, m \geq 0$

$$|\tilde{y}_{t_3+m} - \tilde{y}_{t_3+n}| < \epsilon$$

meaning that $\tilde{y}_t$ is a Cauchy sequence, and thus converges to a finite constant. It follows immediately that y_t converges to a finite, non-zero constant. $\square$

We are now in a position to prove the proposition.

Proof. Note that

$$t^{1-b} \mu_t = t^{1-b} Q \tilde{\mu}_t = Q \left(t^{1-b} \left(\prod_{k=1}^t \left(I - \frac{1}{\lambda_0 + k} (I - \Lambda_k) \right) \right) \tilde{\mu}_0 \right).$$

The j th element of $t^{1-b} \left(\prod_{k=1}^t \left(I - \frac{1}{\lambda_0 + k} (I - \Lambda_k) \right) \right) \tilde{\mu}_0$ is given by

$$t^{1-b} \left(\prod_{k=1}^t r_{j,k} \right) \exp \left(i \sum_{k=1}^t \varphi_{j,k} \right) \tilde{\mu}_{j,0}$$

If $\Lambda_{j,r} = b$, then by Lemma 7 $\lim_{t \rightarrow \infty} t^{1-b} \left(\prod_{k=1}^t r_{j,k} \right) = \kappa_j$ for some finite, positive κ_j . If $\Lambda_{j,r} < b$ then $\lim_{t \rightarrow \infty} t^{1-b} \left(\prod_{k=1}^t r_{j,k} \right) = 0$. The conclusion of the proposition follows by noting that $\tilde{\mu}_{j*,0} \neq 0$ by assumption, that $\exp(\phi i)$ is bounded for all ϕ , and that the columns of Q are linearly independent. $\square$

A comment related to the possibility of sinusoidal fluctuations in $t^{1-b}\tilde{\mu}_{j^*,t}$ is in order. From the definition of $\varphi_{j^*,k}$ and the power series representation of $\sin^{-1}$

$$\sin(\varphi_{j^*,k}) = \frac{\frac{\Lambda_{j^*,k,c}}{\lambda_0+k}}{r_{j,k}}$$

$$\sin^{-1}(x) = \sum_{n=0}^{\infty} \frac{(2n)!}{(2^n n!)^2} \frac{x^{2n+1}}{2n+1}.$$

By Lemma 6, for large enough k all of the terms in the power series representation of $\sin(\varphi_{j^*,k})$ are of the same sign. It follows that for large enough t it must be that $\sum_{k=1}^t \varphi_{j^*,k}$ is either strictly increasing or strictly decreasing in t and that

$$|\varphi_{j^*,t}| > \frac{\frac{|\Lambda_{j^*,t,c}|}{\lambda_0+k}}{r_{j^*,t}}.$$

By Lemma 6 and because $\lim_{t \rightarrow \infty} r_{j^*,t} = 1$, if $|\Lambda_{j^*,c}| \neq 0$ we have $\lim_{t \rightarrow \infty} \sum_{k=1}^t |\varphi_{j^*,k}| = \infty$. This result, along with the observation that $\lim_{t \rightarrow \infty} \varphi_{j^*,t} = 0$, implies sinusoidal fluctuations in $t^{1-b}\tilde{\mu}_{j^*,t}$.

B Constant gain learning

Another, widely-used way to model learning based on past data is to have people update beliefs using a constant gain. When people update their beliefs using constant-gain learning, equation (4) becomes

$$\mu_t = \mu_{t-1} + \gamma(x_t - \mu_{t-1}),$$

for $0 < \gamma_b < 1$. Rearranging, we obtain the analog of equation (9)

$$\mu_t - \frac{a}{1-b} = \sum_{j=0}^{t-1} (1-\gamma_b)^j \left(\frac{\varepsilon_{t-j}}{1-b} \right) \gamma_b + (1-\gamma_b)^t \left(\mu_0 - \frac{a}{1-b} \right) \quad (46)$$

where $\gamma_b = (1-b)\gamma$. As long as $|1-\gamma_b| < 1$ the impact of μ_0 on μ_t goes to zero eventually. As in the case of Bayesian learning, the rate of convergence is decreasing in b . So, the positive feedback loop discussed in the previous section continues to operate.

Rewriting equation (46)

$$E \frac{\mu_t - \frac{a}{1-b}}{\mu_0 - \frac{a}{1-b}} = (1 - \gamma_b)^t.$$

Here, convergence occurs at a geometric rate, λ^t , $0 \leq \lambda < 1$. In contrast, convergence under Bayesian learning proceeds at a power rate, $t^{-\delta}$, $\delta > 0$ (Proposition 1). Power convergence is well known to be slower than geometric convergence for any $\delta > 0$ and $\lambda < 1$.

Nevertheless, convergence can be very slow under constant gain learning. For example, when $\gamma = 0.5$ and $b = 0, 0.5, 0.75, 0.85, .95$. Again, let T satisfy $(1 - \gamma_b)^T \simeq 1/3$. Then $T = 2(3), 4(11), 9(113), 15(2201)$, and $44(5.2 \text{ billion})$, respectively. Numbers in parentheses reproduce the results under Bayesian learning when $\lambda_0 = 1$. The variable, T , is increasing at an increasing rate as b gets larger. While learning under constant gain learning can be very slow for large b , it is much faster than when the gain is decreasing.

C Solution algorithm for non-linear NK model

In this Appendix we detail our solution strategy for the non-linear NK model we consider in our paper. We exploit the model's structure to simplify its solution. In particular, because the steady state is an absorbing state for the REE and learning equilibria that we consider, we can solve the steady state decision rules without reference to the period when $r = r_\ell$. With this solution in hand, we then turn to the period when $r = r_\ell$, which is where we consider learning.

Our main model code is implemented in c++, with reliance on the Eigen, boost, and nlopt libraries. Our computations were conducted using nearly 400 processors with heavy reliance on MPI. Our computations took roughly two weeks to complete. Details related to our model code are available in the README file associated with the replication materials. This Appendix outlines the strategy used to solve the model that is implemented in that code.

C.1 Steady state

In the steady state, there is no uncertainty. However, households still face a bond-holding choice and firms still face a relative-price choice. In an REE, households will

choose to hold zero bonds and firms will choose to set their price to the aggregate price level.

C.1.1 Household problem

In the steady state, the household value function is given by

$$V_{h,ss}(b_h) = \max_{C_h, N_h, b'_h} \left\{ \log(C_h) - \frac{\chi}{2} (N_h)^2 + \beta V_{h,ss}(b'_h) \right\}$$

subject to

$$C_h + \frac{b'_h}{R_{ss}} \leq \frac{b_h}{\pi_{ss}} + w_{ss} N_h + T_{ss}.$$

Here, b_h and b'_h are household h 's real bond holdings chosen in the previous and current period, respectively. The variables C_h and N_h are household h 's consumption and labor supply. The aggregate variable R_{ss} , π_{ss} , w_{ss} , and T_{ss} are the gross nominal interest rate, the gross inflation rate, the real wage, and taxes net of transfers and profits. The values of these aggregate variables are known to the household. We constrain households so that $b'_h \in [\underline{b}, \bar{b}]$. Implicitly, we have functions $C_{h,ss}(b_h)$, $N_{h,ss}(b_h)$, and $b'_{h,ss}(b_h)$. Assuming the constraint on b'_h is not binding, household maximization implies

$$\frac{1}{C_{h,ss}(b_h)} = \beta R_{ss} \frac{1}{C_{h,ss}(b'_h(b_h)) \pi_{ss}} \quad (47)$$

$$\chi C_{h,ss}(b_h) N_{h,ss}(b_h) = w_{ss} \quad (48)$$

$$C_{h,ss}(b_h) + \frac{b'_{h,ss}(b_h)}{R_{ss}} = \frac{b_h}{\pi_{ss}} + w_{ss} N_{h,ss}(b_h) + T_{ss} \quad (49)$$

We define a grid over $[\underline{b}, \bar{b}]$ and approximate the functions $C_{h,ss}(b_h)$, $N_{h,ss}(b_h)$, and $b'_{h,ss}(b_h)$ on that grid in the following way.³²

- (i) We conjecture a value for $b'_{h,ss}(b_h)$ at each grid point. Call the conjectured value $b'^i_{h,ss}(b_h)$.

³²In our implementation, we set $-\underline{b} = \bar{b} = 1$, which is equal to steady state output. We use a symmetric grid with 25 points that includes zero and places more points near zero than at more extreme values because $b_h = b'_h = 0$ in both REE and in learning equilibria.

(ii) Note that equations (48) and (49) can be written as

$$\chi C_{h,ss}(b_h) \left(C_{h,ss}(b_h) + \frac{b_{h,ss}^i(b_h)}{R_{ss}} - \frac{b_h}{\pi_{ss}} - T_{ss} \right) = w_{ss}^2.$$

The left-hand-side is increasing in $C_{h,ss}(b_h) \geq 0$. For every b_h , we solve for the value of $C_{h,ss}(b_h)$ that makes this hold with equality. We call this the conjectured value for $C_{h,ss}(b_h)$ and denote it by $C_{h,ss}^i(b_h)$. Note that with $C_{h,ss}^i(b_h)$, we can back out $N_{h,ss}^i(b_h)$ from equation (48).

(iii) For each grid point, b_h , find b'_h that solves the following version of equation (47)

$$C_h \beta R_{ss} \frac{1}{C_{h,ss}^i(b'_h)} - 1 = 0$$

where $C_h \geq 0$ solves

$$\chi C_h \left(C_h + \frac{b'_h}{R_{ss}} - \frac{b_h}{\pi_{ss}} - T_{ss} \right) = w_{ss}^2.$$

We use linear interpolation to compute $C_{h,ss}^i(b'_h)$ for values of b'_h that fall between grid points. If the procedure would set $b'_h > \bar{b}$ or $b'_h < \underline{b}$, we set b'_h to the respective endpoint of the grid. We record the value of b'_h in by updating the conjectured rule for $b'_{h,ss}(b_h)$ using $b_{h,ss}^{i+1}(b_h) = b'_h$.

(iv) Having computed $b_{h,ss}^{i+1}(b_h)$ for every grid point, we check to see if

$$|b_{h,ss}^{i+1}(b_h) - b_{h,ss}^i(b_h)| < \epsilon$$

at every grid point for some small ϵ . If yes, we say that we have solved the household problem in steady state. If no, we set $b_{h,ss}^i(b_h) = b_{h,ss}^{i+1}(b_h)$ and repeat steps (ii), (iii), and (iv).

Because $\beta \frac{R_{ss}}{\pi_{ss}} = 1$, it is not surprising that we find that $b'_{h,ss}(b_h) = b_h$.

C.1.2 Firm problem

In the steady state, the firm value function is given by

$$V_{f,ss}(p_f) = \max_{p'_f} \left\{ \frac{1}{C_{ss}} (p'_f - (1 - \nu) w_{ss}) (p'_f)^{-\varepsilon} Y_{ss} \right. \\ \left. - \frac{1}{C_{ss}} \frac{\phi}{2} \left(\frac{p'_f}{p_f} \pi_{ss} - 1 \right)^2 (C_{ss} + G_{ss}) \right. \\ \left. + \beta V_{f,ss}(p'_f) \right\}.$$

Here, p_f and p'_f are the ratio of firm f 's price to the aggregate price level in the previous and current period, respectively. The aggregate values π_{ss} , w_{ss} , C_{ss} , G_{ss} , and Y_{ss} are known to the firm. We constrain firms so that $\log(p'_f) \in [\underline{p}, \bar{p}]$. Implicitly, we have a function $p'_{f,ss}(p_f)$. Assuming the constraint on p'_f is not binding, firm maximization implies

$$\begin{aligned} \phi \left(\frac{p'_f(p_f)}{p_f} \pi_{ss} - 1 \right) \frac{1}{p_f} \pi_{ss} (C_{ss} + G_{ss}) = \\ (\varepsilon - 1) \left(\frac{w_{ss}}{p'_f(p_f)} - 1 \right) (p'_f(p_f))^{-\varepsilon} Y_{ss} \\ + \beta \phi \left(\frac{p'_f(p'_f(p_f))}{p'_f(p_f)} \pi_{ss} - 1 \right) \frac{p'_f(p'_f(p_f))}{(p'_f(p_f))^2} \pi_{ss} (C_{ss} + G_{ss}) \end{aligned} \quad (50)$$

We define a grid over $[\underline{p}, \bar{p}]$ and approximate the function $p'_{f,ss}(p_f)$ on that grid in the following way.³³

- (i) We conjecture a value for $p'_{f,ss}(p_f)$ at each grid point. Call the conjectured value $p'^i_{f,ss}(p_f)$.

³³In our implementation, we set $-\underline{p} = \bar{p} = 1$. We use a symmetric grid with 25 points that includes zero that places more points near zero than at more extreme values because $\log(p_f) = \log(p'_f) = 0$ in both REE and in learning equilibria.

- (ii) For each grid point, p_f , find p'_f that solves the following version of equation (50)

$$\begin{aligned} & \phi \left(\frac{p'_f}{p_f} \pi_{ss} - 1 \right) \frac{1}{p_f} \pi_{ss} (C_{ss} + G_{ss}) = \\ & (\varepsilon - 1) \left(\frac{w_{ss}}{p'_f} - 1 \right) (p'_f)^{-\varepsilon} Y_{ss} \\ & + \beta \phi \left(\frac{p'_{f,ss}(p'_f)}{p'_f} \pi_{ss} - 1 \right) \frac{p'_{f,ss}(p'_f)}{(p'_f)^2} \pi_{ss} (C_{ss} + G_{ss}) \end{aligned}$$

We use linear interpolation over $\log(p'_f)$ to compute $p'_{f,ss}(p'_f)$ for values of $\log(p'_f)$ that fall between grid points. If the procedure would set $\log(p'_f) > \bar{p}$ or $\log(p'_f) < \underline{p}$, we set p'_f to the respective endpoint of the grid. We record the value of p'_f in by updating the conjectured rule for $p'_{f,ss}(p_f)$ using $p'_{f,ss}(p_f) = p'_f$.

- (iii) Having computed $p'_{f,ss}(p_f)$ for every grid point, we check to see if

$$|p'_{f,ss}(p_f) - p'_{f,ss}(p_f)| < \epsilon$$

at every grid point for some small ϵ . If yes, we say that we have solved the firm problem in steady state. If no, we set $p'_{f,ss}(p_f) = p'_{f,ss}(p_f)$ and repeat steps (ii) and (iii).

C.2 Solution when $r = r_\ell$

To address the case when $r = r_\ell$, we assume that we have the steady state decision rules in hand and that households and firms know these decision rules with certainty.

C.2.1 Beliefs

Before presenting the household and firm problems, some comments about beliefs are in order when $r = r_\ell$. To simplify the model, we assume households and firms have the same beliefs (though they do not know that they have the same beliefs). Households and firms believe that so long as $r = r_\ell$ the log of aggregate consumption, $\log(C)$, and the log of aggregate inflation, $\log(\pi)$, have uncorrelated Normal distributions with

unknown means and variances. That is

$$\begin{aligned}\log(\pi) &\sim N(\mu_\pi, \sigma_\pi^2) \\ \log(C) &\sim N(\mu_C, \sigma_C^2).\end{aligned}$$

We assume that households and firms have beliefs about the means and variances of the distributions for $\log(C)$ and $\log(\pi)$ that are characterized by density functions that are proportional to Normal-inverse-gamma distributions. These beliefs are not exactly Normal-inverse-gamma distributions because the households and firms embed in their beliefs an upper bound on the variances. This upper bound is important because if variances were unbounded, $\mathbb{E}[\pi] = \mathbb{E}[C] = \infty$, which would challenge the applicability of an expected utility framework. The distributions characterizing beliefs are independent across C and π . That is, for $i \in \{\pi, C\}$, $\mu_i \in (-\infty, \infty)$ and $\sigma_i^2 \in [0, \bar{\sigma}_i^2]$ we have

$$\begin{aligned}\Pr(\sigma_i^2 | \alpha_i, \beta_i) &= \frac{\frac{\beta_i^{\alpha_i}}{\Gamma(\alpha_i)} \left(\frac{1}{\sigma_i^2}\right)^{\alpha_i+1} \exp\left(-\frac{\beta_i}{\sigma_i^2}\right)}{\frac{\Gamma(\alpha_i, \beta_i/\bar{\sigma}_i^2)}{\Gamma(\alpha_i)}}, \\ \Pr(\mu_i | \sigma_i^2, m_i, \lambda_i) &= \frac{\sqrt{\lambda_i}}{\sqrt{2\pi\sigma_i^2}} \exp\left(-\frac{\lambda_i}{2\sigma_i^2} (\mu_i - m_i)^2\right).\end{aligned}$$

Here, $\Gamma(\cdot)$ is the gamma function and $\Gamma(\cdot, \cdot)$ is the incomplete gamma function. Note that $\Gamma(\cdot) = \Gamma(\cdot, 0)$. Again, the advantage of truncating the support of σ_i^2 is that $\mathbb{E}[\pi] < \infty$ if $\bar{\sigma}_\pi^2 < \infty$ and $\mathbb{E}[C] < \infty$ if $\bar{\sigma}_C^2 < \infty$.

Even though we truncate the distributions for σ_i^2 , we maintain conjugacy between prior and posterior beliefs as well as the usual recursive updating equations because the likelihoods associated with observations of π and C are not truncated. Beliefs about $\log(i)$ are parameterized by four values, α_i , β_i , m_i , and λ_i . So, we have 8 total values for both π and C . The standard recursive updating formulas for these variables are

$$\begin{aligned}\lambda'_i &= \lambda_i + 1 \\ m'_i &= \frac{\lambda m_i + \log(i)}{\lambda + 1} \\ \alpha'_i &= \alpha_i + 1/2 \\ \beta'_i &= \beta_i + \frac{\lambda_i (\log(i) - m_i)^2}{2(\lambda_i + 1)}.\end{aligned}$$

Here, a prime indicates the value taken after having observed $\log(i)$.

We need to include variables in Θ that will fully capture the values α_i , β_i , m_i , and λ_i for $i \in \{\pi, C\}$. First, we keep $\frac{1}{t_\ell}$ in Θ , which is the inverse of the number of periods that r has been equal to r_ℓ . We keep the inverse because it is bounded between zero and one, which will be useful. From this value, we can trivially back out λ_i and α_i , given their values in the first period when $r = r_\ell$. We set the initial value of $\lambda_i = 1$ and the initial value of $\alpha_i = 2$. We keep m_C and m_π in Θ . And we also keep

$$\psi'_i = \sqrt{\psi_i^2 \frac{2\alpha'_i}{2\alpha'_i + 1} + \frac{\lambda_i}{\lambda_i + 1} \frac{1}{2\alpha'_i + 1} (\log(i) - m_i)^2}.$$

Note that by setting $\beta_i = (\psi_i)^2 \alpha'_i$ it is clear that we recover the exactly recursive structure of β_i (given above). An advantage of using ψ_i in Θ rather than β_i is that ψ_i is a consistent estimator for the standard deviation, whereas β_i generally grows without bound (except when the standard deviation is zero). Keeping the values of Θ within bounded grids will be important for the purposes of approximation. In total, $\Theta = \left[\frac{1}{t_\ell}, m_\pi, m_C, \psi_\pi, \psi_C\right]$ has five elements and we have a mapping from Θ to α_i , β_i , m_i , and λ_i for $i \in \{\pi, C\}$. We also have a law of motion for Θ so that $\Theta' = L(\Theta, [\pi, C])$.

An advantage of the Normal-inverse-gamma setup detailed above is that we can have analytic expressions for the distribution for the variables $\log(\pi)$ and $\log(C)$ conditional on Θ . In particular

$$\begin{aligned} \Pr(\log(i) | \Theta) &= \frac{\Pr(\log(i) | \mu_i, \sigma_i^2, \Theta) \Pr(\mu_i, \sigma_i^2 | \Theta)}{\Pr(\mu_i, \sigma_i^2 | \log(i), \Theta)} \\ &= \frac{\frac{1}{\sqrt{2\pi\sigma_i^2}} \exp\left(-\frac{1}{2\sigma_i^2} (\log(i) - m_i)^2\right)}{\frac{\sqrt{\lambda'_i}}{\sqrt{2\pi\sigma_i^2}} \exp\left(-\frac{\lambda'_i}{2\sigma_i^2} (\mu_i - m'_i)^2\right) \frac{(\beta'_i)^{\alpha'_i}}{\Gamma(\alpha'_i)} \left(\frac{1}{\sigma_i^2}\right)^{\alpha'_i+1} \exp\left(-\frac{\beta'_i}{\sigma_i^2}\right)} \\ &\quad \times \frac{\sqrt{\lambda_i}}{\sqrt{2\pi\sigma_i^2}} \exp\left(-\frac{\lambda_i}{2\sigma_i^2} (\mu_i - m_i)^2\right) \\ &\quad \times \frac{\beta_i^{\alpha_i}}{\Gamma(\alpha_i)} \left(\frac{1}{\sigma_i^2}\right)^{\alpha_i+1} \exp\left(-\frac{\beta_i}{\sigma_i^2}\right) \frac{\kappa'_i}{\kappa_i} \end{aligned}$$

where

$$\kappa_i = \frac{\Gamma(\alpha_i, \beta_i/\bar{\sigma}^2)}{\Gamma(\alpha_i)}.$$

Then

$$\begin{aligned} \Pr(\log(i) | \Theta) &= \left(\frac{\lambda_i \alpha_i}{\beta_i (\lambda_i + 1)} \right)^{1/2} \frac{\Gamma(\alpha_i + 1/2)}{\Gamma(\alpha_i) \sqrt{2\pi\alpha_i}} \\ &\times \left(1 + \frac{1}{2\alpha_i} \frac{(\log(i) - m_i)^2}{\left(\frac{\lambda_i \alpha_i}{\beta_i (\lambda_i + 1)} \right)^{-1}} \right)^{-\alpha_i - 1/2} \frac{\kappa'_i}{\kappa_i}. \end{aligned} \quad (51)$$

Notice that κ'_i depends on the point of evaluation for $\log(i)$. Evidently, if we ignored the ratio κ'_i/κ_i , which would be correct in the case when $\bar{\sigma}_i^2 = \infty$, the pdf for $\log(i)$ is a t distribution with location parameter m_i , scale parameter $\left(\frac{\lambda_i \alpha_i}{\beta_i (\lambda_i + 1)} \right)^{-1/2}$, and $2\alpha_i$ degrees of freedom. If $\bar{\sigma}_i^2$ is large, $\kappa'_i/\kappa_i \neq 1$ but is close to unity. For finite $\bar{\sigma}_i^2$, the ratio κ'_i/κ_i serves to thin the tails of the distribution of $\log(i)$ by down-weighting the probability of extreme values for $\log(i)$.³⁴ Because the density function of the t distribution is readily available and reliably computed in statistical software and because κ_i and κ'_i are easily computed using readily available implementations of the gamma and incomplete gamma functions, we can use equation (51) for quadrature weighting. We use Gauss-Hermite quadrature with seven nodes when computing approximations to integrals based on equation (51).

C.2.2 Household problem

When $r = r_\ell$, the household value function is given by

$$\begin{aligned} V_h(b_h, \Theta, x) &= \max_{C_h, N_h, b'_h} \left\{ \log(C_h) - \frac{\chi}{2} (N_h)^2 \right. \\ &\quad \left. + \frac{1}{1 + r_\ell} [p \mathbb{E}_{\Theta'} V_h(b'_h, \Theta', x') + (1 - p) V_{h,ss}(b'_h)] \right\} \end{aligned}$$

subject to

$$C_h + \frac{b'_h}{R} \leq \frac{b_h}{\pi} + wN_h + T.$$

Here, $x = [\pi, C]'$, $V_{h,ss}(\cdot)$ is the steady state value function for the household, which is defined above, and $\mathbb{E}_{\Theta'}$ denotes expectations of the household computed conditional on Θ' . Given x and Θ , we have $\Theta' = L(\Theta, x)$. So, the expectation of the household is

³⁴We set $\bar{\sigma}_i^2$ equal to the squared maximum value on the grid for ψ_i (described below).

taken with respect to x' , which is believed to be iid. We assume that households know the monetary and fiscal policy rules. We also assume that they correctly think that $Y = (C + G) (1 + \frac{\phi}{2} (\pi - 1)^2)$, $N = Y$, and $w = \chi CY$. Given x , with these assumptions R , π , w , and T can be computed. The steady state values of aggregate variables are known to the household. We constrain households so that $b'_h \in [\underline{b}, \bar{b}]$. The household optimization problem gives us implicit functions for $C_h(b_h, \Theta, x)$, $N_h(b_h, \Theta, x)$, and $b'_h(b_h, \Theta, x)$. Considering interior solutions for b'_h , we have

$$\frac{1}{C_h(b_h, \Theta, x)} = \frac{1}{1 + r_\ell} R \left[p \mathbb{E}_{\Theta'} \left\{ \frac{1}{\pi' C'_h(b'_h, \Theta', x')} \right\} + (1 - p) \frac{1}{\pi_{ss} C_{h,ss}(b'_h)} \right] \quad (52)$$

$$w = \chi C_h(b_h, \Theta, x) N_h(b_h, \Theta, x) \quad (53)$$

$$C_h(b_h, \Theta, x) = \frac{b_h}{\pi} + w N_h(b_h, \Theta, x) + T - \frac{b'_h(b_h, \Theta, x)}{R}. \quad (54)$$

Instead of approximating $C_h(b_h, \Theta, x)$, $N_h(b_h, \Theta, x)$, and $b'_h(b_h, \Theta, x)$ directly, we approximate

$$v_h(b_h, \Theta) = \mathbb{E}_\Theta \left\{ \frac{1}{\pi C_h(b_h, \Theta, x)} \right\}.$$

We take this approach because we can eliminate x as a state variable in the approximation. We define grids on the elements of Θ and use the grid defined for b_h in the steady state. We then approximate $v_h(b_h, \Theta)$ in the following way.³⁵

- (i) We conjecture a value for $v_h(b_h, \Theta)$ at each grid point in the cross product of the grids over the elements of b_h and Θ .³⁶ Call the conjectured value $v_h^i(b_h, \Theta)$.
- (ii) For a given grid point we use quadrature to get a value for $\mathbb{E}_\Theta \{(\pi C_h)^{-1}\}$. To solve for the expectation of interest, we need to solve for C_h given many different values for x . Conditional on a value for x , equations (53) and (54) can be written as

$$\chi C_h \left(C_h + \frac{b'_h}{R} - \frac{b_h}{\pi} - T \right) = w^2$$

³⁵The grids for m_i contain 12 points that are not evenly spaced. They include each REE point as well as the target-inflation steady state. The remaining points are bunched relatively close to the REE points. The grid for ψ_C contains 11 points that are evenly spaced from 0 to 0.1. The grid for ψ_π contains 11 points that are evenly spaced from 0 to 0.05. Note that inflation is expressed in quarterly terms, so a change of 0.05 would be 20 percent if annualized.

³⁶There are $435,600 = 12 \times 12 \times 11 \times 11 \times 25$ points in the cross product of the grids for m_i , ψ_i , and b_h . The grid for t_ℓ^{-1} is handled in a way discussed below.

The left-hand-side is increasing in $C_h \geq 0$. For a given b'_h , we solve for the value of C_h that makes this hold with equality. We then search for the value of b'_h that makes the following version of equation (52) hold with equality:

$$\frac{1}{C_h} = \frac{1}{1 + r_\ell} R \left[p v_h^i(b'_h, \Theta') + (1 - p) \frac{1}{\pi_{ss} C_{h,ss}(b'_h)} \right].$$

We use linear interpolation to compute $v_h^i(b'_h, \Theta')$ for values of b'_h and Θ' that fall between grid points. If the procedure would set $b'_h > \bar{b}$ or $b'_h < \underline{b}$, we set b'_h to the respective endpoint of the grid for b_h . We record the associated value of C_h and use it in the quadrature to compute $v_h^{i+1}(b_h, \Theta) = \mathbb{E}_\Theta \{ (\pi C_h)^{-1} \}$.

(iii) Having computed $v_h^{i+1}(b_h, \Theta)$ for every grid point, we check to see if

$$|v_h^i(b_h, \Theta) - v_h^{i+1}(b_h, \Theta)| < \epsilon$$

at every grid point for some small ϵ . If yes, we say that we have solved the household problem when $r = r_\ell$. If no, we set $v_h^i(b_h, \Theta) = v_h^{i+1}(b_h, \Theta)$, repeat steps (ii) and (iii).

The grid that we use on $\frac{1}{t_\ell}$ is special. In particular, we let that grid be $[0, \frac{1}{99}, \frac{1}{98}, \dots, 1]$. The first element of the grid corresponds to the case when infinite time has past. In this case households think that they would update their beliefs so that $\Theta' = \Theta$ because m_i and ψ_i are consistent estimators for the means and variances. In our numerical computations, we utilize this fact to first approximate v_h in this case. We then approximate v_h in the case where $t_\ell = 99$. When $t_\ell = 99$, we need to interpolate between the solution to the case when $t_\ell = \infty$ and the conjectured value of $v_h^i(b_h, \Theta)$ when $t_\ell = 99$ to evaluate $v_h^i(b'_h, \Theta')$. That is, when $t_\ell = 99$ we have to find a fixed point of this interpolation, which is computationally intense. To do the interpolation, we linearly interpolate between $t_\ell^{-1} = 1/99$ and $t_\ell^{-1} = 0$. When $t_\ell = 98$, having approximated $v_h(b_h, \Theta)$ for $t_\ell = 99$ means that can evaluate $v_h(b'_h, \Theta')$ exactly at $t'_\ell = 99$ without reference to $v_h^i(b_h, \Theta)$. We approximate for $v_h(b_h, \Theta)$ when $t_\ell = 98$ and work work back in this way to $t_\ell = 1$. This strategy fits this into the structure of steps 1-3 because we know that the value of $v_h(b_h, \Theta)$ will not depend on its value at any any t_ℓ that is smaller than implied by Θ . So, we have a block dependent structure to $v_h(b_h, \Theta)$. Additionally, we know that t_ℓ will only take integer values.

C.2.3 Firm problem

When $r = r_\ell$, the firm value function is given by

$$V_f(p_f, \Theta, x) = \max_{p'_f} \left\{ \frac{1}{C} \left((p'_f - (1 - \nu) w) (p'_f)^{-\varepsilon} Y - \frac{\phi}{2} \left(\frac{p'_f}{p_f} \pi - 1 \right)^2 (C + G) \right) \right. \\ \left. + \frac{1}{1 + r_\ell} [p \mathbb{E}_{\Theta'} V_f(p'_f, \Theta', x') + (1 - p) V_{f,ss}(p'_f)] \right\}$$

Here, $x = [\pi, C]'$, $V_{f,ss}(\cdot)$ is the steady state value function for the firm, which is defined above, and $\mathbb{E}$ denotes expectations of the firm. Given x and Θ , we have $\Theta' = L(\Theta, x)$. So, the expectation of the firm is taken with respect to x' , which is believed to be iid. We assume that firms know the monetary and fiscal policy rules. We also assume that they correctly think that $Y = (C + G) (1 + \frac{\phi}{2} (\pi - 1)^2)$, $N = Y$, and $w = \chi CY$. Given x , with these assumptions π , w , G , and Y can be computed. The steady state values of aggregate variables are known to the firm. We constrain firms so that $\log(p'_f) \in [\underline{p}, \bar{p}]$. Implicitly, from firm optimization we have a function $p'_f(p_f, \Theta, x)$. Considering interior solutions for p'_f , firm maximization implies

$$\begin{aligned} & \phi \left(\frac{p'_f(p_f, \Theta, x)}{p_f} \pi - 1 \right) \frac{1}{p_f} \pi (C + G) = \\ & (\varepsilon - 1) \left(\frac{w}{p'_f(p_f, \Theta, x)} - 1 \right) (p'_f(p_f, \Theta, x))^{-\varepsilon} Y \\ & + \frac{1}{1 + r_\ell} p \mathbb{E}_{\Theta'} \frac{C}{C'} \phi \left(\frac{p'_f(p'_f(p_f, \Theta, x), \Theta', x'))}{p'_f(p_f, \Theta, x)} \pi' - 1 \right) \frac{p'_f(p'_f(p_f, \Theta, x), \Theta', x'))}{(p'_f(p_f, \Theta, x))^2} \pi' (C' + G') \\ & + \frac{1}{1 + r_\ell} \frac{C}{C_{ss}} (1 - p) \phi \left(\frac{p'_{f,ss}(p'_f(p_f, \Theta, x))}{p'_f(p_f, \Theta, x)} \pi_{ss} - 1 \right) \frac{p'_{f,ss}(p'_f(p_f, \Theta, x))}{(p'_f(p_f, \Theta, x))^2} \pi_{ss} (C_{ss} + G_{ss}). \end{aligned} \tag{55}$$

Instead of approximating $p'_f(p_f, \Theta, x)$ directly, we approximate

$$v_f(p_f, \Theta) = \mathbb{E}_\Theta \left\{ \frac{1}{C} \phi \left(\frac{p'_f}{p_f} \pi - 1 \right) \frac{p'_f}{p_f} \pi (C + G) \right\}.$$

We take this approach because we can eliminate x as a state variable in the approximation. We use the same grids on the elements of Θ that we use for the household problem

and the grid defined for $\log(p_f)$ in the steady state and we approximate $v_f(p_f, \Theta)$ in the following way.

- (i) We conjecture a value for $v_f(p_f, \Theta)$ at each grid point in the cross product of the grids over the elements of p_f and Θ . Call the conjectured value $v_f^i(p_f, \Theta)$.
- (ii) For a given grid point we use quadrature to get a value for

$$\mathbb{E} \left\{ \frac{1}{C} \phi \left(\frac{p'_f}{p_f} \pi - 1 \right) \frac{p'_f}{p_f} \pi (C + G) \right\}.$$

To solve for the expectation of interest, we need to solve for p'_f given many different values for x . Conditional on a value for x , we find a value of p'_f that solves the following version of equation (55)

$$\begin{aligned} & \phi \left(\frac{p'_f}{p_f} \pi - 1 \right) \frac{1}{p_f} \pi (C + G) = \\ & (\varepsilon - 1) \left(\frac{w}{p'_f} - 1 \right) (p'_f)^{-\varepsilon} Y \\ & + \frac{1}{1 + r_\ell} p v_f^i(p'_f, \Theta') \frac{C}{p'_f} \\ & + \frac{1}{1 + r_\ell} \frac{C}{C_{ss}} (1 - p) \phi \left(\frac{p'_{f,ss}(p'_f)}{p'_f} \pi_{ss} - 1 \right) \frac{p'_{f,ss}(p'_f)}{(p'_f)^2} \pi_{ss} (C_{ss} + G_{ss}). \end{aligned}$$

We use linear interpolation over $\log(p'_f)$ to compute $v_f^i(p'_f, \Theta')$ for values of $\log(p'_f)$ and Θ' that fall between grid points. If the procedure would set $\log(p'_f) > \bar{p}$ or $\log(p'_f) < \underline{p}$, we set p'_f to the respective endpoint of the grid for p_f . We record the value of p'_f in and the associated aggregate variables so that the quadrature procedure can approximate

$$v_f^{i+1}(p_f, \Theta) = \mathbb{E}_\Theta \left\{ \frac{1}{C} \phi \left(\frac{p'_f}{p_f} \pi - 1 \right) \frac{p'_f}{p_f} \pi (C + G) \right\}.$$

- (iii) Having computed $v_f^{i+1}(p_f, \Theta)$ for every grid point, we check to see if

$$|v_f^i(p_f, \Theta) - v_f^{i+1}(p_f, \Theta)| < \epsilon$$

at every grid point for some small ϵ . If yes, we say that we have solved the

household problem when $r = r_\ell$. If no, we set $v_f^i(p_f, \Theta) = v_f^{i+1}(p_f, \Theta)$, repeat steps (ii) and (iii).

Our use of the same grids as in the household problem allows us to exploit the same block dependent structure in t_ℓ^{-1} .

C.3 Learning equilibria

Here we detail how we construct learning equilibria, given the solutions to the household and firm problems— v_h and v_f .

(i) Set $r = r_\ell$ and assume a value for Θ_t for $t = 1$.

(ii) Conjecture a value for π_t .

(a) Find the value of C_t that would make the following equation hold

$$\frac{1}{C_t} = \frac{1}{1 + r_\ell} R_t \left[p v_h(0, f(\Theta_t, [\pi_t, C_t])) + (1 - p) \frac{1}{\pi_{ss} C_{h,ss}(0)} \right].$$

Note that with π_1 and C_1 the values of all other aggregate variables can be computed.

(b) Check to see if the following equation holds

$$\begin{aligned} \phi(\pi_t - 1) \pi_t (C_t + G_t) = \\ (\varepsilon - 1)(w_t - 1) + \frac{1}{1 + r_\ell} p v_f(1, f(\Theta_t, [\pi_t, C_t])) C_t \\ + \frac{1}{1 + r_\ell} \frac{C_t}{C_{ss}} (1 - p) \phi(\pi_{ss} - 1) \pi_{ss} (C_{ss} + G_{ss}). \end{aligned}$$

If yes, we have a period equilibrium for period t and we record π_t and C_t . If no, conjecture a different value for π_t .

(iii) Set $\Theta_{t+1} = L(\Theta_t, [\pi_t, C_t])$ and repeat step (ii).

When we consider “anticipated utility,” we define $\tilde{\Theta}_t$ to be Θ_t , but with $\frac{1}{t_\ell} = 0$. We then perform step 2 with $\tilde{\Theta}_t$ instead of Θ_t . However, in step 3 we continue to use Θ_t . The switch between $\tilde{\Theta}_t$ and Θ_t highlights the way in which “anticipated utility” is not internally rational.

D Linearized NK Model

Here we describe our strategy for linearizing the NK model around an REE. We find it convenient to use t notation, rather than recursive notation.

D.1 Household problem

The household have a flow budget constraint

$$C_{h,t} + \frac{b_{h,t}}{R_t} = \frac{b_{h,t-1}}{\pi_t} + w_t N_{h,t} + \tau_t$$

and optimality conditions given by

$$\begin{aligned} \frac{1}{C_{h,t}} \frac{1}{R_t} &= \beta_t \mathbb{E}_{h,t} \frac{1}{C_{h,t+1} \pi_{t+1}} \\ \chi N_{h,t} C_{h,t} &= w_t. \end{aligned}$$

Here, $\mathbb{E}_{h,t}$ is $\mathbb{E}_{\Theta'}$ in our recursive notation. We assume that β_t takes two values: $\tilde{\beta} = \frac{1}{1+r_\ell}$ and β , with $\tilde{\beta} > \beta$. The high value happens at period 1 and goes back to the low value with probability $1 - p$. The low value is the absorbing state.

Let's first consider the absorbing state. Log-linearize (except for $b_{h,t}$, which is linearized) the equilibrium conditions around the zero inflation steady state (note that the aggregate variables take their steady state value and the households know this, so their log-deviation is zero).

$$\begin{aligned} C \hat{C}_{h,t} + \beta \hat{b}_{h,t} &= \hat{b}_{h,t-1} + \hat{N}_{h,t} \\ \hat{C}_{h,t} &= \mathbb{E}_{h,t} [\hat{C}_{h,t+1}] \\ 0 &= \hat{N}_{h,t} + \hat{C}_{h,t} \end{aligned}$$

Evidently,

$$\begin{aligned} \hat{C}_{h,t} &= \frac{1}{C+1} \hat{b}_{h,t-1} - \frac{\beta}{C+1} \hat{b}_{h,t} \\ \hat{C}_{h,t} &= \mathbb{E}_{h,t} [\hat{C}_{h,t+1}] \\ \frac{\beta}{C+1} \hat{b}_{h,t+1} &= \frac{1+\beta}{C+1} \hat{b}_{h,t} - \frac{1}{C+1} \hat{b}_{h,t-1} \end{aligned}$$

meaning

$$\frac{1}{C+1}\widehat{b}_{h,t-1} - \frac{\beta+1}{C+1}\widehat{b}_{h,t} = -\frac{\beta}{C+1}\mathbb{E}_{h,t}\left[\widehat{b}_{h,t+1}\right]$$

We consider solutions of the form

$$\widehat{b}_{h,t} = \omega_{b,b}\widehat{b}_{h,t-1}$$

where $\omega_{b,b}$ satisfies

$$\frac{1}{C+1} - \frac{\beta+1}{C+1}\omega_{b,b} + \frac{\beta}{C+1}\omega_{b,b}^2 = 0.$$

The solutions to this equation are

$$\omega_{b,b} = \frac{\frac{\beta+1}{C+1} \pm \sqrt{\frac{\beta+1}{C+1}^2 - 4\frac{1}{C+1}\frac{\beta}{C+1}}}{2\frac{\beta}{C+1}}$$

We focus on $\omega_{b,b} = 1$ because that is the value that corresponds to the solution of the nonlinear model. So,

$$\begin{aligned}\widehat{b}_{h,t-1} &= \frac{1}{1-\beta}(C+1)\widehat{C}_{h,t} \\ \widehat{b}_{h,t} &= \widehat{b}_{h,t-1}.\end{aligned}$$

Let's next consider the case where $\beta_t = \tilde{\beta}$. Let $\tilde{x}$ be the RE aggregate quantity while $\beta_t = \tilde{\beta}$ and $\hat{x}_t$ be the (log-)linearized quantity around $\tilde{x}$. We have

$$\begin{aligned}\tilde{C}\widehat{C}_{h,t} + \widehat{b}_{h,t} &= \frac{\widehat{b}_{h,t-1}}{\tilde{\pi}} + \tilde{w}\tilde{N}\widehat{w}_t + \tilde{w}\tilde{N}\widehat{N}_{h,t} + \tilde{\tau}\widehat{\tau}_t \\ -\frac{1}{\tilde{C}\tilde{R}}\widehat{C}_{h,t} &= -\tilde{\beta}\left[\frac{p}{\tilde{C}\tilde{\pi}}\mathbb{E}_{h,t}\left(\widehat{C}_{h,t+1} + \widehat{\pi}_{t+1}\right) + \frac{(1-p)}{C\pi}\frac{1-\beta}{C+1}\widehat{b}_{h,t}\right] \\ \widehat{w}_t &= \widehat{N}_{h,t} + \widehat{C}_{h,t}\end{aligned}$$

Note that we have imposed $R_t = 1$ while $\beta_t = \tilde{\beta}$, which is true in the REE. In this

sense, the system is local. We assume that households know that

$$\begin{aligned} N_t &= (C_t + G) \left(1 + \frac{\Phi}{2} (\pi_t - 1)^2 \right) \\ w_t &= \chi N_t C_t \\ \tau_t &= (1 - w_t) Y_t - \frac{\Phi}{2} (\pi_t - 1)^2 (C_t + G) - G. \end{aligned}$$

These relations are true in the period equilibrium, and are log-linearized to be

$$\begin{aligned} \hat{N}_t &= \left(1 + \frac{\Phi}{2} (\tilde{\pi} - 1)^2 \right) \frac{\tilde{C}}{\tilde{N}} \hat{C}_t + \Phi \left(\frac{\tilde{C} + \tilde{G}}{\tilde{N}} \right) (\tilde{\pi} - 1) \tilde{\pi} \hat{\pi}_t \\ \hat{w}_t &= \hat{N}_t + \hat{C}_t \\ \tilde{\tau} \hat{\tau}_t &= (1 - \tilde{w}) \tilde{Y} \hat{Y}_t - \tilde{w} \tilde{Y} \hat{w}_t - \frac{\Phi}{2} (\tilde{\pi} - 1)^2 \tilde{C} \hat{C}_t \\ &\quad - \Phi (\tilde{\pi} - 1) \tilde{\pi} (\tilde{C} + \tilde{G}) \hat{\pi}_t. \end{aligned}$$

The household optimality conditions and aggregate relations that are known to the household can be written as a single equation of the form =

$$\begin{aligned} & -\kappa_{b,t-1} \hat{b}_{h,t-1} + \kappa_{b,t} \hat{b}_{h,t} - \kappa_{C,t} \hat{C}_t - \kappa_{\pi,t} \hat{\pi}_t \\ & = \kappa_{b,t+1} \mathbb{E}_{h,t} \left(\hat{b}_{h,t+1} \right) - \kappa_{\mu_C,t} m_{C,t} - \kappa_{\mu_\pi,t} m_{\pi,t} \end{aligned}$$

where

$$\begin{aligned}
\kappa_{b,t-1} &= \frac{1}{\tilde{\pi}} \\
\kappa_{b,t} &= 1 + \tilde{\beta} \tilde{R} \left(\frac{p}{\tilde{\pi}^2} + \tilde{C} \left(\tilde{C} + \tilde{w} \tilde{N} \right) \frac{1-p}{C\pi} \frac{1-\beta}{C+1} \right) \\
\kappa_{C,t} &= 2\tilde{w} \tilde{N} \left(\left(1 + \frac{\Phi}{2} (\tilde{\pi} - 1)^2 \right) \frac{\tilde{C}}{\tilde{N}} + 1 \right) \\
&\quad - (2\tilde{w} - 1) \left(1 + \frac{\Phi}{2} (\tilde{\pi} - 1)^2 \right) \tilde{C} \\
&\quad - \left(\tilde{w} \tilde{Y} + \frac{\Phi}{2} (\tilde{\pi} - 1)^2 \tilde{C} \right) \\
\kappa_{\pi,t} &= 0 \\
\kappa_{b,t+1} &= \tilde{\beta} \tilde{R} \frac{p}{\tilde{\pi}} \\
\kappa_{\mu_C,t} &= \tilde{\beta} \tilde{R} \frac{p}{\tilde{\pi}} \kappa_{C,t} \\
\kappa_{\mu_\pi,t} &= \tilde{\beta} \tilde{R} \left(\tilde{C} + \tilde{w} \tilde{N} \right) \frac{p}{\tilde{\pi}}
\end{aligned}$$

Here, $m_{\pi,t}$ is m'_π in our recursive notation and $m_{C,t}$ is m'_C in our recursive notation. Note that the time subscripts on the κ 's is to denote if the coefficient multiplies, for example, b_t or b_{t-1} . The time subscript does not indicate time-variation in the coefficient.

We consider solutions to this equation of the form

$$\hat{b}_{h,t} = \gamma_{b,b} \hat{b}_{h,t-1} + \gamma_{b,\pi} \hat{\pi}_t + \gamma_{b,C} \hat{C}_t + \gamma_{b,\mu_\pi} m_{\pi,t} + \gamma_{b,\mu_C} m_{C,t}.$$

Note that this is a linear approximation to $b'_h(b_h, \Theta)$. Our approximation does not include ψ_π or ψ_C because of the certainty equivalence of the linearized model.

Using the linear decision rule for $\hat{b}_{h,t}$, $\gamma_{b,b}$ is determined by

$$-\kappa_{b,t-1} + \kappa_{b,t} \gamma_{b,b} = \kappa_{b,t+1} \gamma_{b,b}^2$$

which is given by

$$\gamma_{b,b} = \frac{\kappa_{b,t} \pm \sqrt{\kappa_{b,t}^2 - 4\kappa_{b,t-1}\kappa_{b,t+1}}}{2\kappa_{b,t+1}}$$

Both of the solutions for $\gamma_{b,b}$ are larger than unity. However, the smaller value is closer to the solution of the nonlinear model at the REE, so we focus on that value. We

determine the other four values of $\gamma_{b,i}$ using the following equations.

$$\begin{aligned}
\kappa_{b,t}\gamma_{b,\pi} - \kappa_{\pi,t} &= \kappa_{b,t+1}\gamma_{b,b}\gamma_{b,\pi} \\
\kappa_{b,t}\gamma_{b,C} - \kappa_{C,t} &= \kappa_{b,t+1}\gamma_{b,b}\gamma_{b,C} \\
\kappa_{b,t}\gamma_{b,\mu_\pi} &= \kappa_{b,t+1}\gamma_{b,b}\gamma_{b,\mu_\pi} + \kappa_{b,t+1}(\gamma_{b,\mu_\pi} + \gamma_{b,\pi}) - \kappa_{\mu_\pi,t} \\
\kappa_{b,t}\gamma_{b,\mu_C} &= \kappa_{b,t+1}\gamma_{b,b}\gamma_{b,\mu_C} + \kappa_{b,t+1}(\gamma_{b,\mu_C} + \gamma_{b,C}) - \kappa_{\mu_C,t}.
\end{aligned}$$

The first two equations imply

$$\begin{aligned}
\gamma_{b,\pi} &= \frac{\kappa_{\pi,t}}{\kappa_{b,t} - \kappa_{b,t+1}\gamma_{b,b}} \\
\gamma_{b,C} &= \frac{\kappa_{C,t}}{\kappa_{b,t} - \kappa_{b,t+1}\gamma_{b,b}}.
\end{aligned}$$

Then the third and fourth equations imply

$$\begin{aligned}
\gamma_{b,\mu_\pi} &= \frac{\kappa_{\mu_\pi,t} - \kappa_{b,t+1}\gamma_{b,\pi}}{\kappa_{b,t+1}(\gamma_{b,b} + 1) - \kappa_{b,t}} \\
\gamma_{b,\mu_C} &= \frac{\kappa_{\mu_C,t} - \kappa_{b,t+1}\gamma_{b,C}}{\kappa_{b,t+1}(\gamma_{b,b} + 1) - \kappa_{b,t}}
\end{aligned}$$

This gives a solution to the household problem.

D.2 Household problem ignoring the ZLB

We wanted to know what would happen if we ignored the ZLB. In that case, the nominal interest rate is set so that

$$\tilde{R}_t = \frac{1}{\beta} + \alpha(\tilde{\pi}_t - 1) \Rightarrow \hat{\tilde{R}}_t = \alpha \frac{\hat{\tilde{\pi}}}{\tilde{R}} \hat{\tilde{\pi}}_t$$

Then $\kappa_{\pi,t}$ becomes

$$\kappa_{\pi,t} = (\tilde{C} + \tilde{w}\tilde{N}) \tilde{R}^{-1} \alpha \tilde{\pi}$$

and the rest of the analysis in the previous sub-section goes through.

D.3 Firm problem

The firm's optimality condition is

$$\begin{aligned} (p_{f,t} - w_t) (p_{f,t})^{-\varepsilon} Y_t + \frac{\Phi}{\varepsilon - 1} \left(\frac{p_{f,t}}{p_{f,t-1}} \pi_t - 1 \right) \frac{p_{f,t}}{p_{f,t-1}} \pi_t (C_t + G_t) \\ = \beta_t \mathbb{E}_{f,t} \frac{C_t}{C_{t+1}} \frac{\Phi}{\varepsilon - 1} \left(\frac{p_{f,t+1}}{p_{f,t}} \pi_{t+1} - 1 \right) \frac{p_{f,t+1}}{p_{f,t}} \pi_{t+1} (C_{t+1} + G_{t+1}) \end{aligned}$$

We log-linearize this condition for the case when $\beta_t = \beta$ and firms know the steady state values of the variables to get

$$\widehat{p}_{f,t} + \frac{\Phi}{\varepsilon - 1} (\widehat{p}_{f,t} - \widehat{p}_{f,t-1}) = \beta \frac{\Phi}{\varepsilon - 1} (\widehat{p}_{f,t+1} - \widehat{p}_{f,t})$$

We assume a solution of the form

$$\widehat{p}_{f,t} = \omega_{p,p} \widehat{p}_{t-1}$$

so that

$$0 = \beta \frac{\Phi}{\varepsilon - 1} \omega_{p,p}^2 - \left(1 + (1 + \beta) \frac{\Phi}{\varepsilon - 1} \right) \omega_{p,p} + \frac{\Phi}{\varepsilon - 1}$$

which has solutions

$$\omega_{p,p} = \frac{\left(1 + (1 + \beta) \frac{\Phi}{\varepsilon - 1} \right) \pm \sqrt{\left(1 + (1 + \beta) \frac{\Phi}{\varepsilon - 1} \right)^2 - 4\beta \left(\frac{\Phi}{\varepsilon - 1} \right)^2}}{2\beta \frac{\Phi}{\varepsilon - 1}}.$$

Only one of these solutions is less than 1 in absolute value and we use that solution because it resembles our non-linear solution.

Now we will consider the case when $\beta_t = \tilde{\beta}$. In this case,

$$\begin{aligned} (\tilde{p}_{f,t} - \tilde{w}_t) (\tilde{p}_{f,t})^{-\varepsilon} \tilde{Y}_t + \frac{\Phi}{\varepsilon - 1} \left(\frac{\tilde{p}_{f,t}}{\tilde{p}_{f,t-1}} \tilde{\pi}_t - 1 \right) \frac{\tilde{p}_{f,t}}{\tilde{p}_{f,t-1}} \tilde{\pi}_t (\tilde{C}_t + \tilde{G}_t) = \\ p \tilde{\beta} \mathbb{E}_{f,t} \frac{\tilde{C}_t}{\tilde{C}_{t+1}} \frac{\Phi}{\varepsilon - 1} \left(\frac{\tilde{p}_{f,t+1}}{\tilde{p}_{f,t}} \tilde{\pi}_{t+1} - 1 \right) \frac{\tilde{p}_{f,t+1}}{\tilde{p}_{f,t}} \tilde{\pi}_{t+1} (\tilde{C}_{t+1} + \tilde{G}_{t+1}) \\ + (1 - p) \tilde{\beta} \mathbb{E}_{f,t} \frac{\tilde{C}_t}{C_{t+1}} \frac{\Phi}{\varepsilon - 1} \left(\frac{p_{f,t+1}}{\tilde{p}_{f,t}} \pi_{t+1} - 1 \right) \frac{p_{f,t+1}}{\tilde{p}_{f,t}} \pi_{t+1} (C_{t+1} + G_{t+1}) \end{aligned}$$

We log-linearize this to be

$$\begin{aligned}
& \tilde{Y} (1 + \varepsilon (\tilde{w} - 1)) \hat{\tilde{p}}_{f,t} + (1 - \tilde{w}) \tilde{Y} \hat{\tilde{Y}}_t - \tilde{w} \tilde{Y} \hat{\tilde{w}}_t \\
& + \frac{\Phi}{\varepsilon - 1} \tilde{\pi} (\tilde{C} + \tilde{G}) (2\tilde{\pi} - 1) (\hat{\tilde{p}}_{f,t} - \hat{\tilde{p}}_{f,t-1} + \hat{\tilde{\pi}}_t) + \frac{\Phi}{\varepsilon - 1} (\tilde{\pi} - 1) \tilde{\pi} (\tilde{C} \hat{\tilde{C}}_t + \tilde{G} \hat{\tilde{G}}_t) = \\
& p\tilde{\beta} \frac{\Phi}{\varepsilon - 1} (\tilde{\pi} - 1) \tilde{\pi} (\tilde{C} + \tilde{G}) (\hat{\tilde{C}}_t - \mathbb{E}_{f,t} \hat{\tilde{C}}_{t+1}) + \\
& p\tilde{\beta} \frac{\Phi}{\varepsilon - 1} \tilde{\pi} (\tilde{C} + \tilde{G}) (2\tilde{\pi} - 1) (\mathbb{E}_{f,t} \hat{\tilde{p}}_{f,t+1} - \hat{\tilde{p}}_{f,t} + \hat{\tilde{\pi}}_{t+1}) \\
& + p\tilde{\beta} \frac{\Phi}{\varepsilon - 1} (\tilde{\pi} - 1) \tilde{\pi} (\tilde{C} \mathbb{E}_{f,t} \hat{\tilde{C}}_{t+1} + \tilde{G} \mathbb{E}_{f,t} \hat{\tilde{G}}_{t+1}) + \\
& (1 - p) \tilde{\beta} \frac{\tilde{C}}{C} \frac{\Phi}{\varepsilon - 1} (C + G) (\mathbb{E}_{f,t} \hat{\tilde{p}}_{f,t+1} - \hat{\tilde{p}}_{f,t} + \hat{\tilde{\pi}}_{t+1})
\end{aligned}$$

Using

$$\begin{aligned}
\hat{\tilde{Y}}_t &= \hat{\tilde{N}}_t = \left(1 + \frac{\Phi}{2} (\tilde{\pi} - 1)^2\right) \frac{\tilde{C}}{\tilde{N}} \hat{\tilde{C}}_t + \Phi \left(\frac{\tilde{C} + \tilde{G}}{\tilde{N}}\right) (\tilde{\pi} - 1) \tilde{\pi} \hat{\tilde{\pi}}_t \\
\hat{\tilde{w}}_t &= \hat{\tilde{N}}_t + \hat{\tilde{C}}_t
\end{aligned}$$

we can write the firm's optimality condition as

$$\begin{aligned}
& \zeta_{p_f,t} \hat{\tilde{p}}_{f,t} + \zeta_{\pi,t} \hat{\tilde{\pi}}_t - \zeta_{p_f,t-1} \hat{\tilde{p}}_{f,t-1} = \\
& \zeta_{C,t} \hat{\tilde{C}}_t - \zeta_{\mu_C,t} m_{C,t} + \zeta_{p_f,t+1} \mathbb{E}_{f,t} \hat{\tilde{p}}_{f,t+1} + \zeta_{\mu_\pi,t} m_{\pi,t}
\end{aligned}$$

wherex

$$\begin{aligned}
\zeta_{p_f,t} &= \tilde{Y} (1 + \varepsilon (\tilde{w} - 1)) + \frac{\Phi}{\varepsilon - 1} \tilde{\pi} (\tilde{C} + \tilde{G}) (2\tilde{\pi} - 1) \\
&\quad + p\tilde{\beta} \frac{\Phi}{\varepsilon - 1} \tilde{\pi} (\tilde{C} + \tilde{G}) (2\tilde{\pi} - 1) \\
&\quad + (1 - p) \tilde{\beta} \frac{\tilde{C}}{C} \frac{\Phi}{\varepsilon - 1} \pi (C + G) (1 - \omega_{pp}) \\
\zeta_{p_f,t-1} &= \frac{\Phi}{\varepsilon - 1} \tilde{\pi} (\tilde{C} + \tilde{G}) (2\tilde{\pi} - 1) \\
\zeta_{p_f,t+1} &= p\tilde{\beta} \frac{\Phi}{\varepsilon - 1} \tilde{\pi} (\tilde{C} + \tilde{G}) (2\tilde{\pi} - 1) \\
\zeta_{\pi,t} &= \frac{\Phi}{\varepsilon - 1} (\tilde{C} + \tilde{G}) (2\tilde{\pi} - 1) \tilde{\pi} + (1 - 2\tilde{w}) \Phi (\tilde{C} + \tilde{G}) (\tilde{\pi} - 1) \tilde{\pi} \\
\zeta_{C,t} &= - (1 - 2\tilde{w}) \left(1 + \frac{\Phi}{2} (\tilde{\pi} - 1)^2 \right) \tilde{C} + \tilde{w} \tilde{Y} \\
&\quad - \frac{\Phi}{\varepsilon - 1} (\tilde{\pi} - 1) \tilde{\pi} \tilde{C} + p\tilde{\beta} \frac{\Phi}{\varepsilon - 1} (\tilde{\pi} - 1) \tilde{\pi} (\tilde{C} + \tilde{G}) \\
\zeta_{\mu_C,t} &= p\tilde{\beta} \frac{\Phi}{\varepsilon - 1} (\tilde{\pi} - 1) \tilde{\pi} \tilde{G} \\
\zeta_{\mu_\pi,t} &= p\tilde{\beta} \frac{\Phi}{\varepsilon - 1} \tilde{\pi} (\tilde{C} + \tilde{G}) (2\tilde{\pi} - 1).
\end{aligned}$$

As in the household problem, the time subscripts on the ζ 's is to denote if the coefficient multiplies, for example, $p_{f,t}$ or $p_{f,t-1}$. The time subscript does not indicate time-variation in the coefficient. For similar reasons to the solution to the household problem, we consider solutions to this equation of the form

$$\hat{p}_{f,t} = \gamma_{p,p} \hat{p}_{f,t-1} + \gamma_{p,\pi} \hat{\pi}_t + \gamma_{p,C} \hat{C}_t + \gamma_{p,\mu_\pi} m_{\pi,t} + \gamma_{p,\mu_C} m_{C,t}.$$

Note that $\gamma_{p,p}$ is determined by

$$\zeta_{p_f,t+1} \gamma_{p,p}^2 - \zeta_{p_f,t} \gamma_{p,p} + \zeta_{p_f,t-1} = 0$$

So,

$$\gamma_{p,p} = \frac{\zeta_{p_f,t} \pm \sqrt{\zeta_{p_f,t}^2 - 4\zeta_{p_f,t+1}\zeta_{p_f,t-1}}}{2\zeta_{p_f,t+1}}.$$

The smaller root (which is stable) is a better approximation like the nonlinear model.

We then have that

$$\begin{aligned}
\zeta_{p_f,t}\gamma_{p,\pi} + \zeta_{\pi,t} &= \zeta_{p_f,t+1}\gamma_{p,p}\gamma_{p,\pi} \\
\zeta_{p_f,t}\gamma_{p,C} - \zeta_{C,t} &= \zeta_{p_f,t+1}\gamma_{p,p}\gamma_{p,C} \\
\zeta_{p_f,t}\gamma_{p,\mu_\pi} - \zeta_{\mu_\pi,t} &= \zeta_{p_f,t+1}\gamma_{p,p}\gamma_{p,\mu_\pi} + \zeta_{p_f,t+1}(\gamma_{p,\pi} + \gamma_{p,\mu_\pi}) \\
\zeta_{p_f,t}\gamma_{p,\mu_C} + \zeta_{\mu_C,t} &= \zeta_{p_f,t+1}\gamma_{p,p}\gamma_{p,\mu_C} + \zeta_{p_f,t+1}(\gamma_{p,C} + \gamma_{p,\mu_C})
\end{aligned}$$

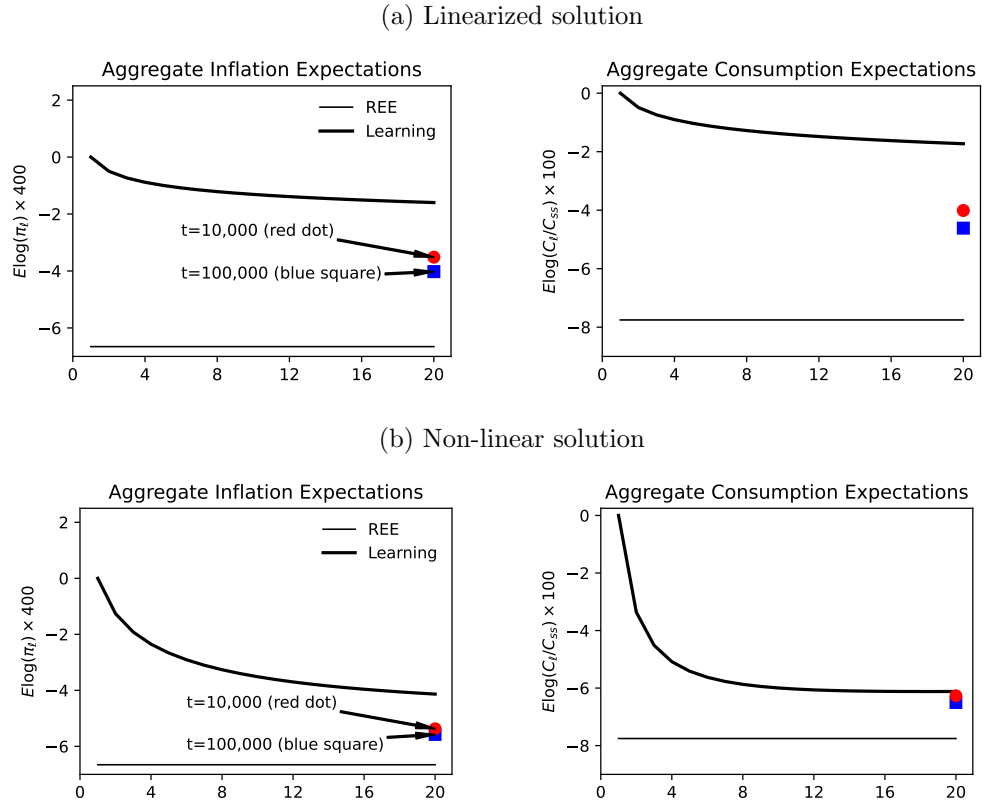
So,

$$\begin{aligned}
\gamma_{p,\pi} &= -\frac{\zeta_{\pi,t}}{\zeta_{p_f,t} - \zeta_{p_f,t+1}\gamma_{p,p}} \\
\gamma_{p,C} &= \frac{\zeta_{C,t}}{\zeta_{p_f,t} - \zeta_{p_f,t+1}\gamma_{p,p}} \\
\gamma_{p,\mu_\pi} &= \frac{\zeta_{\mu_\pi,t} + \zeta_{p_f,t+1}\gamma_{p,\pi}}{\zeta_{p_f,t} - \zeta_{p_f,t+1}\gamma_{p,p} - \zeta_{p_f,t+1}} \\
\gamma_{p,\mu_C} &= \frac{-\zeta_{\mu_C,t} + \zeta_{p_f,t+1}\gamma_{p,C}}{\zeta_{p_f,t} - \zeta_{p_f,t+1}\gamma_{p,p} - \zeta_{p_f,t+1}}.
\end{aligned}$$

Because R_t does not enter the firm optimality condition, ignoring the ZLB has no effect on the linearization of the firm problem.

D.4 Slow convergence in the linearized solution

Figure 10: Slow convergence of beliefs is similar in linearized and non-linear solutions



Note: In the sub-figures (a) and (b) m_i is initially set to the steady state REE value. In all sub-figures, $\psi_i = 0.02$, $\lambda_i = 1$, $\alpha_i = 2$. Source: Authors' calculations.