

Slow Learning

Lawrence Christiano

based on joint work with Martin Eichenbaum and Ben Johannsen

May 28, 2024

Motivation

- In the past two decades, several events with little precedence occurred:
 - ▶ US Financial Crisis, European Sovereign Debt crisis, Covid, Ukraine, Climate Crisis.

Motivation

- In the past two decades, several events with little precedence occurred:
 - ▶ US Financial Crisis, European Sovereign Debt crisis, Covid, Ukraine, Climate Crisis.
- Our standard models assume rational expectations (RE)

Motivation

- In the past two decades, several events with little precedence occurred:
 - ▶ US Financial Crisis, European Sovereign Debt crisis, Covid, Ukraine, Climate Crisis.
- Our standard models assume rational expectations (RE)
 - ▶ Assumes people know *a lot* about the economy:
 - ★ what can happen, the associated probabilities, etc.

Motivation

- In the past two decades, several events with little precedence occurred:
 - ▶ US Financial Crisis, European Sovereign Debt crisis, Covid, Ukraine, Climate Crisis.
- Our standard models assume rational expectations (RE)
 - ▶ Assumes people know *a lot* about the economy:
 - ★ what can happen, the associated probabilities, etc.
 - ▶ Maybe Looked OK during the Great Moderation.

Motivation

- In the past two decades, several events with little precedence occurred:
 - ▶ US Financial Crisis, European Sovereign Debt crisis, Covid, Ukraine, Climate Crisis.
- Our standard models assume rational expectations (RE)
 - ▶ Assumes people know *a lot* about the economy:
 - ★ what can happen, the associated probabilities, etc.
 - ▶ Maybe Looked OK during the Great Moderation.
 - ▶ Harder to justify in unprecedented situations.

What we Do

- Consider situation in which people don't have Rational Expectations and instead learn from observations as time passes.

What we Do

- Consider situation in which people don't have Rational Expectations and instead learn from observations as time passes.
 - ▶ For REE to be useful for policy analysis, require fast convergence to REE.

What we Do

- Consider situation in which people don't have Rational Expectations and instead learn from observations as time passes.
 - ▶ For REE to be useful for policy analysis, require fast convergence to REE.
- Ask: What features of the economy determine speed of convergence to REE?

What we Do

- Consider situation in which people don't have Rational Expectations and instead learn from observations as time passes.
 - ▶ For REE to be useful for policy analysis, require fast convergence to REE.
- Ask: What features of the economy determine speed of convergence to REE?
 - ▶ Use a reduced form example which suggests a simple *learning principle*:
 - ★ When expectations of a variable are partially self-fulfilling, then learning converges *slowly* to REE, if at all.

What we Do

- Consider situation in which people don't have Rational Expectations and instead learn from observations as time passes.
 - ▶ For REE to be useful for policy analysis, require fast convergence to REE.
- Ask: What features of the economy determine speed of convergence to REE?
 - ▶ Use a reduced form example which suggests a simple *learning principle*:
 - ★ When expectations of a variable are partially self-fulfilling, then learning converges *slowly* to REE, if at all.
- Turn to a particular 'event without precedence':
 - ▶ The drop in R to its zero lower bound (ZLB) in 2009-2015.

What we Do

- Consider situation in which people don't have Rational Expectations and instead learn from observations as time passes.
 - ▶ For REE to be useful for policy analysis, require fast convergence to REE.
- Ask: What features of the economy determine speed of convergence to REE?
 - ▶ Use a reduced form example which suggests a simple *learning principle*:
 - ★ When expectations of a variable are partially self-fulfilling, then learning converges *slowly* to REE, if at all.
- Turn to a particular 'event without precedence':
 - ▶ The drop in R to its zero lower bound (ZLB) in 2009-2015.
- Ask: Is convergence fast enough for REE to be a useful laboratory in the ZLB?

What we Do

- Consider situation in which people don't have Rational Expectations and instead learn from observations as time passes.
 - ▶ For REE to be useful for policy analysis, require fast convergence to REE.
- Ask: What features of the economy determine speed of convergence to REE?
 - ▶ Use a reduced form example which suggests a simple *learning principle*:
 - ★ When expectations of a variable are partially self-fulfilling, then learning converges *slowly* to REE, if at all.
- Turn to a particular 'event without precedence':
 - ▶ The drop in R to its zero lower bound (ZLB) in 2009-2015.
- Ask: Is convergence fast enough for REE to be a useful laboratory in the ZLB?
 - ▶ Answer: No.

What we Do

- Consider situation in which people don't have Rational Expectations and instead learn from observations as time passes.
 - ▶ For REE to be useful for policy analysis, require fast convergence to REE.
- Ask: What features of the economy determine speed of convergence to REE?
 - ▶ Use a reduced form example which suggests a simple *learning principle*:
 - ★ When expectations of a variable are partially self-fulfilling, then learning converges *slowly* to REE, if at all.
- Turn to a particular 'event without precedence':
 - ▶ The drop in R to its zero lower bound (ZLB) in 2009-2015.
- Ask: Is convergence fast enough for REE to be a useful laboratory in the ZLB?
 - ▶ Answer: No.
 - ▶ For the classic NK model, convergence is *extremely* slow in the ZLB.

What we Do

- Consider situation in which people don't have Rational Expectations and instead learn from observations as time passes.
 - ▶ For REE to be useful for policy analysis, require fast convergence to REE.
- Ask: What features of the economy determine speed of convergence to REE?
 - ▶ Use a reduced form example which suggests a simple *learning principle*:
 - ★ When expectations of a variable are partially self-fulfilling, then learning converges *slowly* to REE, if at all.
- Turn to a particular 'event without precedence':
 - ▶ The drop in R to its zero lower bound (ZLB) in 2009-2015.
- Ask: Is convergence fast enough for REE to be a useful laboratory in the ZLB?
 - ▶ Answer: No.
 - ▶ For the classic NK model, convergence is *extremely* slow in the ZLB.
 - ★ The ZLB would almost certainly be over long before learning has come close to converging.

What we Do

- Consider situation in which people don't have Rational Expectations and instead learn from observations as time passes.
 - ▶ For REE to be useful for policy analysis, require fast convergence to REE.
- Ask: What features of the economy determine speed of convergence to REE?
 - ▶ Use a reduced form example which suggests a simple *learning principle*:
 - ★ When expectations of a variable are partially self-fulfilling, then learning converges *slowly* to REE, if at all.
- Turn to a particular 'event without precedence':
 - ▶ The drop in R to its zero lower bound (ZLB) in 2009-2015.
- Ask: Is convergence fast enough for REE to be a useful laboratory in the ZLB?
 - ▶ Answer: No.
 - ▶ For the classic NK model, convergence is *extremely* slow in the ZLB.
 - ★ The ZLB would almost certainly be over long before learning has come close to converging.
 - ▶ Relate this result to the learning principle.

Outline

- Simple example:
 - ▶ Learning principle.
- New Keynesian analysis of shocks and policies in the ZLB using nonlinear version Eggertsson-Woodford (2003) model.
 - ▶ Government spending
 - ▶ Forward Guidance
 - ▶ Interpret results using learning principle.

Simple Example: REE

- Model analyzed in Bray and Savin (ECMA1986):

$$x_t = a + b\mathbb{E}_{t-1}x_t + \varepsilon_t, \varepsilon_t \sim iin(0, \sigma^2), \sigma^2 < \infty$$

‘Workhorse model’ for learning (see, e.g., Evans and Honkapohja (2001)).

► structures

Simple Example: REE

- Model analyzed in Bray and Savin (ECMA1986):

$$x_t = a + b\mathbb{E}_{t-1}x_t + \varepsilon_t, \varepsilon_t \sim iin(0, \sigma^2), \sigma^2 < \infty$$

'Workhorse model' for learning (see, e.g., Evans and Honkapohja (2001)).

► structures

- We consider the following parameter values: $-\infty < b < 1$
 - ▶ When $b < 0$: Muth's (1961) version of Cobweb model,
 - ▶ when $b > 0$, Lucas (1973) 'aggregate supply model'

Simple Example: REE

- Model analyzed in Bray and Savin (ECMA1986):

$$x_t = a + b\mathbb{E}_{t-1}x_t + \varepsilon_t, \varepsilon_t \sim iiN(0, \sigma^2), \sigma^2 < \infty$$

'Workhorse model' for learning (see, e.g., Evans and Honkapohja (2001)). ► structures

- We consider the following parameter values: $-\infty < b < 1$
 - ▶ When $b < 0$: Muth's (1961) version of Cobweb model,
 - ▶ when $b > 0$, Lucas (1973) 'aggregate supply model'
- Rational expectations equilibrium:

$$\mathbb{E}_{t-1}x_t = E_{t-1}x_t, x_t = \overbrace{\frac{a}{1-b}}^{\mu} + \varepsilon_t.$$

- In REE, $x_t \sim iiN(\mu, \sigma^2)$.

Simple Example: Learning

- Bayesian Learning about μ (assume people know the form of the REE process and value of σ^2)

Simple Example: Learning

- Bayesian Learning about μ (assume people know the form of the REE process and value of σ^2)
 - ▶ In period 0, prior on μ is $N\left(\mu_0, \frac{\sigma^2}{\lambda_0}\right)$, where $\lambda_0 \geq 0$ is a measure of precision of prior.

Simple Example: Learning

- Bayesian Learning about μ (assume people know the form of the REE process and value of σ^2)
 - ▶ In period 0, prior on μ is $N\left(\mu_0, \frac{\sigma^2}{\lambda_0}\right)$, where $\lambda_0 \geq 0$ is a measure of precision of prior.
 - ▶ In period t observe $x_1, \dots, x_t$, so Bayes' rule implies posterior $N\left(\mu_t, \frac{\sigma^2}{\lambda_0 + t}\right)$ and

$$\mu_t = \mu_{t-1} + \frac{1}{\lambda_0 + t} (x_t - \mu_{t-1})$$

$$x_t = a + b\mu_{t-1} + \varepsilon_t$$

▶

Simple Example: Learning

- Bayesian Learning about μ (assume people know the form of the REE process and value of σ^2)
 - ▶ In period 0, prior on μ is $N\left(\mu_0, \frac{\sigma^2}{\lambda_0}\right)$, where $\lambda_0 \geq 0$ is a measure of precision of prior.
 - ▶ In period t observe $x_1, \dots, x_t$, so Bayes' rule implies posterior $N\left(\mu_t, \frac{\sigma^2}{\lambda_0 + t}\right)$ and

$$\mu_t = \mu_{t-1} + \frac{1}{\lambda_0 + t} (x_t - \mu_{t-1})$$

$$x_t = a + b\mu_{t-1} + \varepsilon_t$$

- ▶ How people learn is a fundamental part of the law of motion of the system.

Simple Example: Learning

- Bayesian Learning about μ (assume people know the form of the REE process and value of σ^2)
 - ▶ In period 0, prior on μ is $N\left(\mu_0, \frac{\sigma^2}{\lambda_0}\right)$, where $\lambda_0 \geq 0$ is a measure of precision of prior.
 - ▶ In period t observe $x_1, \dots, x_t$, so Bayes' rule implies posterior $N\left(\mu_t, \frac{\sigma^2}{\lambda_0 + t}\right)$ and

$$\mu_t = \mu_{t-1} + \frac{1}{\lambda_0 + t} (x_t - \mu_{t-1})$$

$$x_t = a + b\mu_{t-1} + \varepsilon_t$$

- ▶ How people learn is a fundamental part of the law of motion of the system.
- Repeated substitution:

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

where

$$z_t = \prod_{j=1}^t (1 - b_j), \quad b_j = \frac{1-b}{\lambda_0 + j}.$$

Simple Example: Convergence Questions

- does $\mu_t \rightarrow \mu = a/(1 - b)$?

Simple Example: Convergence Questions

- does $\mu_t \rightarrow \mu = a/(1 - b)$?
 - ▶ Yes for $b < 1$.
 - ▶ This result is known at least since Bray and Savin (1986).

Simple Example: Convergence Questions

- does $\mu_t \rightarrow \mu = a/(1 - b)$?
 - ▶ Yes for $b < 1$.
 - ▶ This result is known at least since Bray and Savin (1986).
- how fast does convergence occur?

Simple Example: Convergence Questions

- does $\mu_t \rightarrow \mu = a/(1 - b)$?
 - ▶ Yes for $b < 1$.
 - ▶ This result is known at least since Bray and Savin (1986).
- how fast does convergence occur?
 - ▶ potentially, very slowly.

A Feedback Loop and Speed of Convergence

- To understand convergence rate, recall data-generating process under learning:

$$x_t = a + b\mu_{t-1} + \varepsilon_t$$

$$\mu_t = \mu_{t-1} + \frac{1}{\lambda_0 + t} (x_t - \mu_{t-1})$$



A Feedback Loop and Speed of Convergence

- To understand convergence rate, recall data-generating process under learning:

$$x_t = a + b\mu_{t-1} + \varepsilon_t$$

$$\mu_t = \mu_{t-1} + \frac{1}{\lambda_0 + t} (x_t - \mu_{t-1})$$

- ▶ There is a *feedback loop* $\mu_{t-1} \rightarrow x_t \rightarrow \mu_t \rightarrow x_{t+1} \dots$

A Feedback Loop and Speed of Convergence

- To understand convergence rate, recall data-generating process under learning:

$$x_t = a + b\mu_{t-1} + \varepsilon_t$$

$$\mu_t = \mu_{t-1} + \frac{1}{\lambda_0 + t} (x_t - \mu_{t-1})$$

- ▶ There is a *feedback loop* $\mu_{t-1} \rightarrow x_t \rightarrow \mu_t \rightarrow x_{t+1} \dots$
- ▶ If $1 > b > 0$: feedback loop is positive and expectations are (partially) self-fulfilling.

A Feedback Loop and Speed of Convergence

- To understand convergence rate, recall data-generating process under learning:

$$x_t = a + b\mu_{t-1} + \varepsilon_t$$

$$\mu_t = \mu_{t-1} + \frac{1}{\lambda_0 + t} (x_t - \mu_{t-1})$$

- ▶ There is a *feedback loop* $\mu_{t-1} \rightarrow x_t \rightarrow \mu_t \rightarrow x_{t+1} \dots$
- ▶ If $1 > b > 0$: feedback loop is positive and expectations are (partially) self-fulfilling.
 - ★ People slow to leave their initial prior, μ_0 .

A Feedback Loop and Speed of Convergence

- To understand convergence rate, recall data-generating process under learning:

$$x_t = a + b\mu_{t-1} + \varepsilon_t$$

$$\mu_t = \mu_{t-1} + \frac{1}{\lambda_0 + t} (x_t - \mu_{t-1})$$

- ▶ There is a *feedback loop* $\mu_{t-1} \rightarrow x_t \rightarrow \mu_t \rightarrow x_{t+1} \dots$
- ▶ If $1 > b > 0$: feedback loop is positive and expectations are (partially) self-fulfilling.
 - ★ People slow to leave their initial prior, μ_0 .
- ▶ If $b < 0$ expectations self-defeating.

A Feedback Loop and Speed of Convergence

- To understand convergence rate, recall data-generating process under learning:

$$x_t = a + b\mu_{t-1} + \varepsilon_t$$

$$\mu_t = \mu_{t-1} + \frac{1}{\lambda_0 + t} (x_t - \mu_{t-1})$$

- ▶ There is a *feedback loop* $\mu_{t-1} \rightarrow x_t \rightarrow \mu_t \rightarrow x_{t+1} \dots$
- ▶ If $1 > b > 0$: feedback loop is positive and expectations are (partially) self-fulfilling.
 - ★ People slow to leave their initial prior, μ_0 .
- ▶ If $b < 0$ expectations self-defeating.
 - ★ People may be quick to shift away from μ_0 .

A Feedback Loop and Speed of Convergence

- To understand convergence rate, recall data-generating process under learning:

$$x_t = a + b\mu_{t-1} + \varepsilon_t$$
$$\mu_t = \mu_{t-1} + \frac{1}{\lambda_0 + t} (x_t - \mu_{t-1})$$

- ▶ There is a *feedback loop* $\mu_{t-1} \rightarrow x_t \rightarrow \mu_t \rightarrow x_{t+1} \dots$
- ▶ If $1 > b > 0$: feedback loop is positive and expectations are (partially) self-fulfilling.
 - ★ People slow to leave their initial prior, μ_0 .
- ▶ If $b < 0$ expectations self-defeating.
 - ★ People may be quick to shift away from μ_0 .
- Suggests speed of convergence may be a *decreasing* function of b .

Simple Example: Learning Might be Very Slow (or, Fast)

- Consider expected gap relative to REE, as fraction of initial gap:

$$z_t = \frac{E\left(\mu_t - \frac{a}{1-b}\right)}{\mu_0 - \frac{a}{1-b}} = f(t, \lambda_0, b).$$

How long does it take to close 2/3 of initial gap, $z_T = 1/3$?

Simple Example: Learning Might be Very Slow (or, Fast)

- Consider expected gap relative to REE, as fraction of initial gap:

$$z_t = \frac{E\left(\mu_t - \frac{a}{1-b}\right)}{\mu_0 - \frac{a}{1-b}} = f(t, \lambda_0, b).$$

How long does it take to close 2/3 of initial gap, $z_T = 1/3$?

- Answer ($\lambda_0 = 1$):

b	0	0.5	0.75	0.85	0.95
T					

Simple Example: Learning Might be Very Slow (or, Fast)

- Consider expected gap relative to REE, as fraction of initial gap:

$$z_t = \frac{E\left(\mu_t - \frac{a}{1-b}\right)}{\mu_0 - \frac{a}{1-b}} = f(t, \lambda_0, b).$$

How long does it take to close 2/3 of initial gap, $z_T = 1/3$?

- Answer ($\lambda_0 = 1$):

b	0	0.5	0.75	0.85	0.95
T	2				

Simple Example: Learning Might be Very Slow (or, Fast)

- Consider expected gap relative to REE, as fraction of initial gap:

$$z_t = \frac{E\left(\mu_t - \frac{a}{1-b}\right)}{\mu_0 - \frac{a}{1-b}} = f(t, \lambda_0, b).$$

How long does it take to close 2/3 of initial gap, $z_T = 1/3$?

- Answer ($\lambda_0 = 1$):

b	0	0.5	0.75	0.85	0.95
T	2	10			

Simple Example: Learning Might be Very Slow (or, Fast)

- Consider expected gap relative to REE, as fraction of initial gap:

$$z_t = \frac{E\left(\mu_t - \frac{a}{1-b}\right)}{\mu_0 - \frac{a}{1-b}} = f(t, \lambda_0, b).$$

How long does it take to close 2/3 of initial gap, $z_T = 1/3$?

- Answer ($\lambda_0 = 1$):

b	0	0.5	0.75	0.85	0.95
T	2	10	120		

Simple Example: Learning Might be Very Slow (or, Fast)

- Consider expected gap relative to REE, as fraction of initial gap:

$$z_t = \frac{E\left(\mu_t - \frac{a}{1-b}\right)}{\mu_0 - \frac{a}{1-b}} = f(t, \lambda_0, b).$$

How long does it take to close 2/3 of initial gap, $z_T = 1/3$?

- Answer ($\lambda_0 = 1$):

b	0	0.5	0.75	0.85	0.95
T	2	10	120	2500	

Simple Example: Learning Might be Very Slow (or, Fast)

- Consider expected gap relative to REE, as fraction of initial gap:

$$z_t = \frac{E\left(\mu_t - \frac{a}{1-b}\right)}{\mu_0 - \frac{a}{1-b}} = f(t, \lambda_0, b).$$

How long does it take to close 2/3 of initial gap, $z_T = 1/3$?

- Answer ($\lambda_0 = 1$):

b	0	0.5	0.75	0.85	0.95
T	2	10	120	2500	4 billion

Simple Example: Learning Might be Very Slow (or, Fast)

- Christopeit and Massmann (2018) derive various asymptotic properties, μ_t , as $t \rightarrow \infty$.
 - ▶ For example, $z_t \simeq \kappa t^{b-1}$, $\kappa \neq 0$ as $t \rightarrow \infty$, for $b < 1$.

Simple Example: Learning Might be Very Slow (or, Fast)

- Christopeit and Massmann (2018) derive various asymptotic properties, μ_t , as $t \rightarrow \infty$.
 - ▶ For example, $z_t \simeq \kappa t^{b-1}$, $\kappa \neq 0$ as $t \rightarrow \infty$, for $b < 1$.
- Learning principle:
 - ▶ positive feedback loop ($b > 0$): slow learning.
 - ▶ negative feedback loop ($b < 0$): relatively fast learning. ▶ constant gain

Turning to New Keynesian Model

- Recursive Formulation of NK Model

Turning to New Keynesian Model

- Recursive Formulation of NK Model
- Key findings:
 - ▶ When the ZLB model is binding, NK model corresponds to a *high- b economy*.

Turning to New Keynesian Model

- Recursive Formulation of NK Model
- Key findings:
 - ▶ When the ZLB model is binding, NK model corresponds to a *high- b economy*.
 - ★ Model in ZLB implies a strong positive feedback loop in inflation expectations.

Turning to New Keynesian Model

- Recursive Formulation of NK Model
- Key findings:
 - ▶ When the ZLB model is binding, NK model corresponds to a *high- b economy*.
 - ★ Model in ZLB implies a strong positive feedback loop in inflation expectations.
 - ▶ Convergence to a REE is very slow.

Turning to New Keynesian Model

- Recursive Formulation of NK Model
- Key findings:
 - ▶ When the ZLB model is binding, NK model corresponds to a *high- b economy*.
 - ★ Model in ZLB implies a strong positive feedback loop in inflation expectations.
 - ▶ Convergence to a REE is very slow.
 - ▶ REE analysis very misleading for government spending multiplier and forward guidance.

NK Model with Learning

- Simple closed economy, NK model without capital, flexible wages, Rotemberg-sticky prices.
 - ▶ Up to period 0, economy is in unique steady state REE with
 - ★ $\beta = 1/(1 + r_{ss})$, $ss \sim$ 'steady state'
 - ★ gross nominal interest rate, $R > 1$.

NK Model with Learning

- Simple closed economy, NK model without capital, flexible wages, Rotemberg-sticky prices.
 - ▶ Up to period 0, economy is in unique steady state REE with
 - ★ $\beta = 1/(1 + r_{ss})$, $ss \sim$ 'steady state'
 - ★ gross nominal interest rate, $R > 1$.
- In period 0, everyone discovers unexpectedly that r drops to $r_\ell < r_{ss}$ (Eggertsson-Woodford, 2003).
 - ▶ People know the law of motion of r , $r \in (r_\ell, r_{ss})$, r_{ss} is an absorbing state and $P[r_{t+1} = r_\ell | r_t = r_\ell] = p$.
 - ▶ When economy reverts to absorbing state, $r = r_{ss}$, everyone understands we're back to unique steady state REE with $R > 1$.

Model

- What people in the model don't know:
 - ▶ how the economy will evolve over time during the ZLB.
 - ▶ the dynamic general equilibrium effects of government policies.

Model

- What people in the model don't know:
 - ▶ how the economy will evolve over time during the ZLB.
 - ▶ the dynamic general equilibrium effects of government policies.
- People learn about these things as data come in.
 - ▶ Circular process: learning influenced by the data and data influenced by learning.

Model

- What people in the model don't know:
 - ▶ how the economy will evolve over time during the ZLB.
 - ▶ the dynamic general equilibrium effects of government policies.
- People learn about these things as data come in.
 - ▶ Circular process: learning influenced by the data and data influenced by learning.
- Learning:
 - ▶ When data arrive, they update beliefs using Bayes' rule.
 - ▶ In thinking about the future they internalize that their beliefs will continue evolving as new data arrive.

Households

- Beginning of Period State Variables for h^{th} household, $h \in (0, 1)$:
 - ▶ $b_h \sim$ stock of bonds acquired in previous period.

Households

- Beginning of Period State Variables for h^{th} household, $h \in (0, 1)$:
 - ▶ $b_h \sim$ stock of bonds acquired in previous period.
 - ▶ $r \sim$ discount rate observed at the beginning of the period.

Households

- Beginning of Period State Variables for h^{th} household, $h \in (0, 1)$:
 - ▶ $b_h \sim$ stock of bonds acquired in previous period.
 - ▶ $r \sim$ discount rate observed at the beginning of the period.
 - ▶ $\Theta \sim$ parameters governing beliefs about density of x .

Households

- Beginning of Period State Variables for h^{th} household, $h \in (0, 1)$:
 - ▶ $b_h \sim$ stock of bonds acquired in previous period.
 - ▶ $r \sim$ discount rate observed at the beginning of the period.
 - ▶ $\Theta \sim$ parameters governing beliefs about density of x .
 - ★ $x = [C, \pi] \sim$ aggregate variables that allow people to deduce R (nominal interest rate), w (real wage), T (profits net of lump sum taxes)

Households

- Beginning of Period State Variables for h^{th} household, $h \in (0, 1)$:
 - ▶ $b_h \sim$ stock of bonds acquired in previous period.
 - ▶ $r \sim$ discount rate observed at the beginning of the period.
 - ▶ $\Theta \sim$ parameters governing beliefs about density of x .
 - ★ $x = [C, \pi] \sim$ aggregate variables that allow people to deduce R (nominal interest rate), w (real wage), T (profits net of lump sum taxes)
 - ★ Density of x degenerate when $r = r_{ss}$, non-trivial with $r = r_\ell$.

Households

- Beginning of Period State Variables for h^{th} household, $h \in (0, 1)$:
 - ▶ $b_h \sim$ stock of bonds acquired in previous period.
 - ▶ $r \sim$ discount rate observed at the beginning of the period.
 - ▶ $\Theta \sim$ parameters governing beliefs about density of x .
 - ★ $x = [C, \pi] \sim$ aggregate variables that allow people to deduce R (nominal interest rate), w (real wage), T (profits net of lump sum taxes)
 - ★ Density of x degenerate when $r = r_{ss}$, non-trivial with $r = r_\ell$.
- The h^{th} household forms plans for C_h, N_h, b'_h contingent on the not-yet-realized current value of x .

Household x -Contingent Plan

- For a range of values of $x = [C, \pi]$ the h^{th} household chooses C_h, N_h, b'_h to solve:

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

Household x -Contingent Plan

- For a range of values of $x = [C, \pi]$ the h^{th} household chooses C_h, N_h, b'_h to solve:

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

subject to the budget constraint:

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

where V_h and V_h^{ss} denote the value functions in case $r = r^\ell$ or $r = r^{ss}$ in the next period, respectively. [► EquilibriumFunction](#)

Household x -Contingent Plan

- For a range of values of $x = [C, \pi]$ the h^{th} household chooses C_h, N_h, b'_h to solve:

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

subject to the budget constraint:

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

where V_h and V_h^{ss} denote the value functions in case $r = r^\ell$ or $r = r^{ss}$ in the next period, respectively. [► EquilibriumFunction](#)

- Here,
 - $\mathbb{E}$ denotes the expectation operator over marginal data density of x' , conditional on $r' = r_\ell, \Theta, x$.

Household x -Contingent Plan

- For a range of values of $x = [C, \pi]$ the h^{th} household chooses C_h, N_h, b'_h to solve:

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

subject to the budget constraint:

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

where V_h and V_h^{ss} denote the value functions in case $r = r^\ell$ or $r = r^{ss}$ in the next period, respectively. [▶ EquilibriumFunction](#)

- Here,
 - ▶ $\mathbb{E}$ denotes the expectation operator over marginal data density of x' , conditional on $r' = r_\ell, \Theta, x$.
 - ▶ Θ' , next period's belief parameters constructed by combining Θ, x .

Beliefs, Θ

- Because they see the same aggregate data, firms and households have same beliefs about the distribution of $x = [C, \pi]$.

Beliefs, Θ

- Because they see the same aggregate data, firms and households have same beliefs about the distribution of $x = [C, \pi]$.
- People think that both elements of $\log x$ are independently drawn from a different Normal distribution.

Beliefs, Θ

- Because they see the same aggregate data, firms and households have same beliefs about the distribution of $x = [C, \pi]$.
- People think that both elements of $\log x$ are independently drawn from a different Normal distribution.
 - ▶ They are uncertain about the mean and variance of each Normal.
 - ▶ Their joint prior over the means and variances of C and π are (truncated) Normal inverse Wishart.

Beliefs, Θ

- Because they see the same aggregate data, firms and households have same beliefs about the distribution of $x = [C, \pi]$.
- People think that both elements of $\log x$ are independently drawn from a different Normal distribution.
 - ▶ They are uncertain about the mean and variance of each Normal.
 - ▶ Their joint prior over the means and variances of C and π are (truncated) Normal inverse Wishart.
- The vector Θ denotes the parameters that characterize these prior distributions.

Evolution of Beliefs over Time: Internalized Learning

- In making their x -contingent decisions, people internalize that Θ' is a function of Θ and the observed value of x :

$$\Theta' = f(\Theta, x).$$

Here, f has an analytic representation.

Evolution of Beliefs over Time: Internalized Learning

- In making their x -contingent decisions, people internalize that Θ' is a function of Θ and the observed value of x :

$$\Theta' = f(\Theta, x).$$

Here, f has an analytic representation.

- The people in our model are 'internally rational' in the sense of Adam and Marcet 2011.

Evolution of Beliefs over Time: Internalized Learning

- In making their x -contingent decisions, people internalize that Θ' is a function of Θ and the observed value of x :

$$\Theta' = f(\Theta, x).$$

Here, f has an analytic representation.

- The people in our model are 'internally rational' in the sense of Adam and Marcet 2011.
- In period 0, Θ_0 are free parameters.

Household Value Function

- Value function satisfies the following fixed point property:

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

subject to the budget constraint.

Household Value Function

- Value function satisfies the following fixed point property:

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

subject to the budget constraint.

- That households can map from x into the aggregate variables required for their budget constraints corresponds to our assumption that they are good at *static* general equilibrium reasoning.

Household Value Function

- Value function satisfies the following fixed point property:

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

subject to the budget constraint.

- That households can map from x into the aggregate variables required for their budget constraints corresponds to our assumption that they are good at *static* general equilibrium reasoning.
 - ▶ However, they are not good at *dynamic* general equilibrium reasoning.
 - ▶ Their beliefs about the future are distorted.

Production and Firms

- Dixit-Stiglitz formalization standard in NK model.
 - ▶ Final good created by aggregating intermediate goods produced by monopolists.

Production and Firms

- Dixit-Stiglitz formalization standard in NK model.
 - ▶ Final good created by aggregating intermediate goods produced by monopolists.
- Intermediate good firms have sticky prices in the sense of Rotemberg.

Production and Firms

- Dixit-Stiglitz formalization standard in NK model.
 - ▶ Final good created by aggregating intermediate goods produced by monopolists.
- Intermediate good firms have sticky prices in the sense of Rotemberg.
- Intermediate firms' problem expressed in recursive form.

Production and Firms

- Dixit-Stiglitz formalization standard in NK model.
 - ▶ Final good created by aggregating intermediate goods produced by monopolists.
- Intermediate good firms have sticky prices in the sense of Rotemberg.
- Intermediate firms' problem expressed in recursive form.
- Have same beliefs as households.

Government

- Fiscal policy:
 - ▶ Baseline: $G = G_{ss} > 0$ fixed for all r .

Government

- Fiscal policy:
 - ▶ Baseline: $G = G_{ss} > \text{fixed for all } r$.
 - ▶ Alternative: $G = G_\ell > G_{ss}, r = r_\ell, G = G_{ss}, r = r_{ss}$.

Government

- Fiscal policy:

- ▶ Baseline: $G = G_{ss} > \text{fixed for all } r$.
- ▶ Alternative: $G = G_\ell > G_{ss}$, $r = r_\ell$, $G = G_{ss}$, $r = r_{ss}$.
- ▶ Government uses lump sum taxes to balance budget in each period.

Government

- Fiscal policy:

- ▶ Baseline: $G = G_{ss} > \text{fixed for all } r$.
- ▶ Alternative: $G = G_{\ell} > G_{ss}$, $r = r_{\ell}$, $G = G_{ss}$, $r = r_{ss}$.
- ▶ Government uses lump sum taxes to balance budget in each period.

- Monetary policy:

$$R = \max \left\{ 1, \frac{1}{\beta} + \alpha (\pi - 1) \right\}, \alpha > 1$$

- We also consider perturbations on this policy, including forward guidance.

Market Clearing in a *Period Learning Equilibrium*

- Given $r = r_\ell$ and Θ ,
- The vector, $x = [C, \pi]$, is adjusted to ensure goods, bonds and labor markets clear in a way that is consistent with private sector optimization and government policy.
 - ▶ The approach is inspired by [Eusepi, Gibbs and Preston, 2022](#).

Learning Equilibrium

- As long as $r = r_\ell$, economy is a sequence of period learning equilibria.
- When $r = r_{ss}$ economy jumps to $R > 1$ REE steady state.

Is Rational Expectations a Useful Guide for Policy Analysis Under Learning?

- First,
 - ▶ Does learning select one of the (multiple) REE in the ZLB?

Is Rational Expectations a Useful Guide for Policy Analysis Under Learning?

- First,
 - ▶ Does learning select one of the (multiple) REE in the ZLB? ▶ multiple REE
- Second,
 - ▶ How quickly does convergence occur?

Is Rational Expectations a Useful Guide for Policy Analysis Under Learning?

- First,
 - ▶ Does learning select one of the (multiple) REE in the ZLB? ▶ multiple REE
- Second,
 - ▶ How quickly does convergence occur?
- Third,
 - ▶ are predictions of REE about macro stabilization policies robust to learning?

Is Rational Expectations a Useful Guide for Policy Analysis Under Learning?

- First,
 - ▶ Does learning select one of the (multiple) REE in the ZLB? ▶ multiple REE
- Second,
 - ▶ How quickly does convergence occur?
- Third,
 - ▶ are predictions of REE about macro stabilization policies robust to learning?
- Related issue: there are also multiple *steady state* REE's in the NK model (BSGU).
 - ▶ Based on our experiments and the literature, we will focus on the zero inflation steady state.

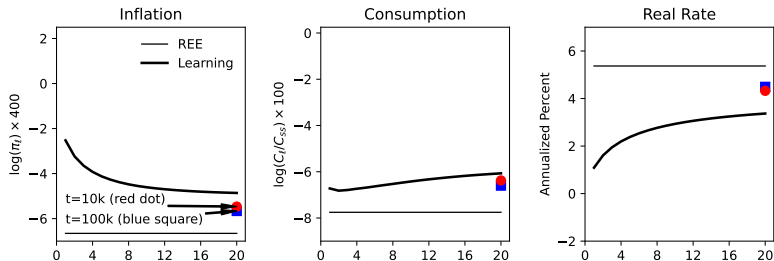
Parameter Values

$$p = 0.80, r_\ell = -0.0015 \text{ } (-0.6APR), G_{ss} = 0.20, r_{ss} = 0.005 \text{ } (2.0APR),$$

$$Y_{ss} = N_{ss} = 1, \varepsilon = 7, \phi = 110, \chi = 1.25, \alpha = 1.5$$

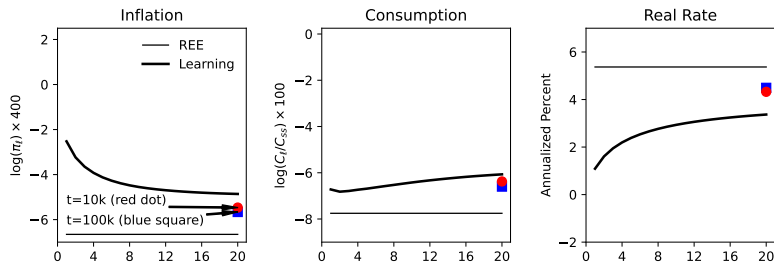
Experiment #1: Slow Learning in the ZLB

- r drops and G remains unchanged.



Experiment #1: Slow Learning in the ZLB

- r drops and G remains unchanged.

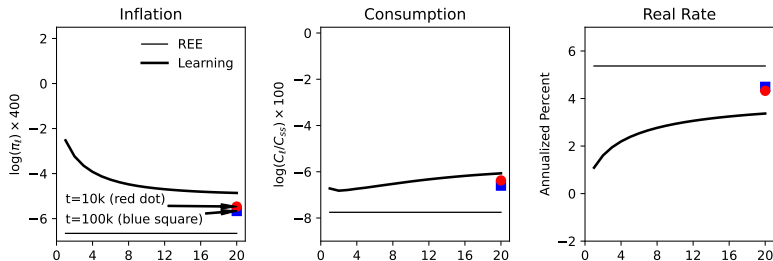


- Key results:

- ▶ Economic impact of the shock under learning is small compared with REE.
- ★ Learning is **extremely** slow.

Experiment #1: Slow Learning in the ZLB

- r drops and G remains unchanged.



- Key results:

- ▶ Economic impact of the shock under learning is small compared with REE.
 - ★ Learning is **extremely** slow.
- ▶ Learning moves the model in the 'right' empirical direction:
 - ★ addresses 'missing deflation puzzle'. ▶ role of internalized learning

Intuition: In ZLB there is a Positive Feedback Loop Between Inflation Expectations and ex post Realized Inflation

Intuition: In ZLB there is a Positive Feedback Loop Between Inflation Expectations and ex post Realized Inflation

- Suppose firms and households *expect lower inflation* in the future during ZLB episode.
 - ▶ Other things the same, firms want to reduce prices now.

Intuition: In ZLB there is a Positive Feedback Loop Between Inflation Expectations and ex post Realized Inflation

- Suppose firms and households *expect lower inflation* in the future during ZLB episode.
 - ▶ Other things the same, firms want to reduce prices now.
 - ▶ Households: $R = 1$ in ZLB, so low inflation expectations $\rightarrow$ real rate high $\rightarrow$ labor supply increased $\rightarrow$ marginal cost of production down $\rightarrow$ inflation down.

Intuition: In ZLB there is a Positive Feedback Loop Between Inflation Expectations and ex post Realized Inflation

- Suppose firms and households *expect lower inflation* in the future during ZLB episode.
 - ▶ Other things the same, firms want to reduce prices now.
 - ▶ Households: $R = 1$ in ZLB, so low inflation expectations $\rightarrow$ real rate high $\rightarrow$ labor supply increased $\rightarrow$ marginal cost of production down $\rightarrow$ inflation down.
 - ▶ With lower actual inflation now, expected inflation in future reduced.

Intuition: In ZLB there is a Positive Feedback Loop Between Inflation Expectations and ex post Realized Inflation

- Suppose firms and households *expect lower inflation* in the future during ZLB episode.
 - ▶ Other things the same, firms want to reduce prices now.
 - ▶ Households: $R = 1$ in ZLB, so low inflation expectations $\rightarrow$ real rate high $\rightarrow$ labor supply increased $\rightarrow$ marginal cost of production down $\rightarrow$ inflation down.
 - ▶ With lower actual inflation now, expected inflation in future reduced.
 - ★ Actual inflation in the future reduced.

Intuition: In ZLB there is a Positive Feedback Loop Between Inflation Expectations and ex post Realized Inflation

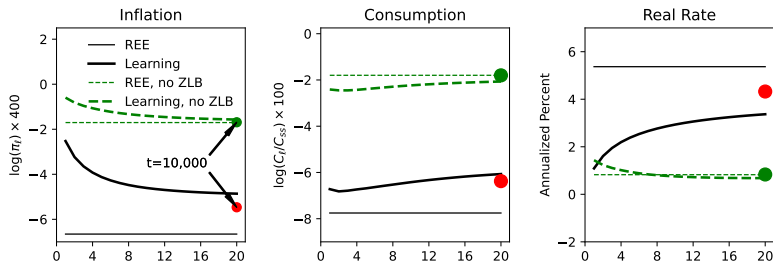
- Suppose firms and households *expect lower inflation* in the future during ZLB episode.
 - ▶ Other things the same, firms want to reduce prices now.
 - ▶ Households: $R = 1$ in ZLB, so low inflation expectations $\rightarrow$ real rate high $\rightarrow$ labor supply increased $\rightarrow$ marginal cost of production down $\rightarrow$ inflation down.
 - ▶ With lower actual inflation now, expected inflation in future reduced.
 - ★ Actual inflation in the future reduced.
- In sum: Households and firms complement each other in creating a positive feedback loop that makes the NK model behave like a 'high- b ' economy.

Role of Binding ZLB in Slow Convergence

- Same shock as above, but lower bound on R ignored.

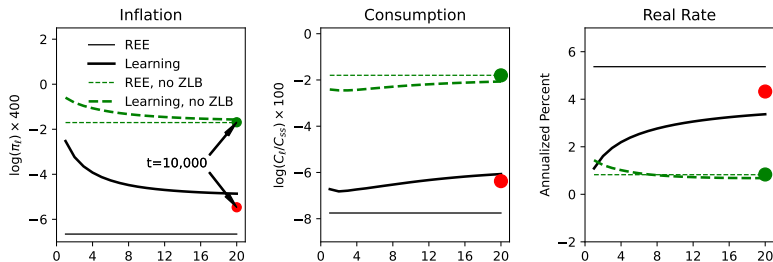
Role of Binding ZLB in Slow Convergence

- Same shock as above, but lower bound on R ignored.



Role of Binding ZLB in Slow Convergence

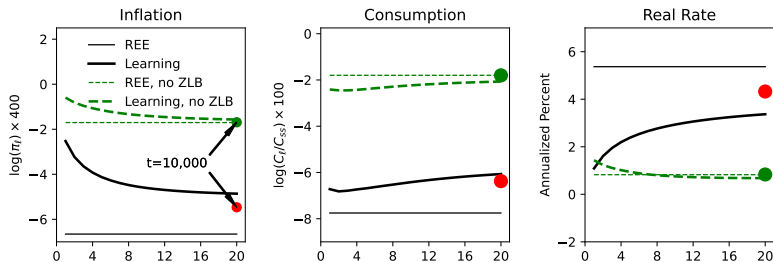
- Same shock as above, but lower bound on R ignored.



- Learning equilibrium converges relatively quickly to REE with Taylor Principle (ZLB ignored).

Role of Binding ZLB in Slow Convergence

- Same shock as above, but lower bound on R ignored.



- Learning equilibrium converges relatively quickly to REE with Taylor Principle (ZLB ignored).
- Reflects learning principle:
 - ▶ Central Bank “Does what it Takes” to keep inflation on target, independent of prior expectations of the public: ‘low b economy’ with Taylor Principle.

Increase in G During ZLB

- Standard result in rational expectations (REE) literature:
 - ▶ multiplier on government spending can be very large in the ZLB.

Increase in G During ZLB

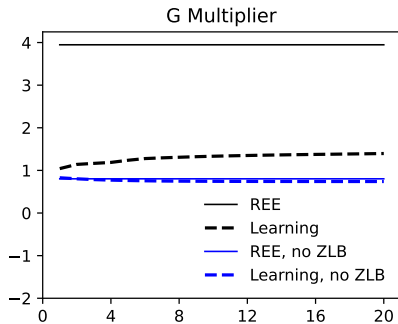
- Standard result in rational expectations (REE) literature:
 - ▶ multiplier on government spending can be very large in the ZLB.
 - ▶ But, large multiplier in REE happens chiefly by raising expected inflation.
 - ★ If learning is backward-looking, then this inflation expectation channel broken.

Increase in G During ZLB

- Standard result in rational expectations (REE) literature:
 - ▶ multiplier on government spending can be very large in the ZLB.
 - ▶ But, large multiplier in REE happens chiefly by raising expected inflation.
 - ★ If learning is backward-looking, then this inflation expectation channel broken.
- Our finding:
 - ▶ We find that the multiplier under learning is very small, compared to REE.
 - ▶ Rational expectations generates *very* misleading prediction about the effects of government spending.

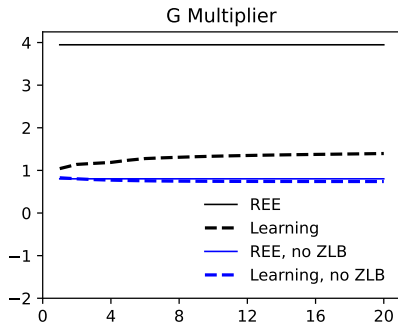
The G Multiplier In ZLB

- Multiplier, $\frac{dY}{dG}$.



The G Multiplier In ZLB

- Multiplier, $\frac{dY}{dG}$.



- A huge difference between REE and learning when Taylor Principle is out of the picture.

We Identify the Analog of b in the Linearized solution to the NK Model in a Neighborhood of a Stable REE Equilibrium

- Let $\hat{\mu}_t = \begin{bmatrix} \hat{\pi}_t \\ \hat{C}_t \end{bmatrix}$ represent the log deviation of people's time t posterior of $\mathbb{E}_t x_{t+1}$ and the REE value of $E x_{t+1}$.
- We show that for sufficiently large t , approximately

$$\hat{x}_t = B \hat{\mu}_{t-1}.$$

- ▶ The structure maps beliefs about $\hat{x}_t$ (i.e., $\hat{\mu}_{t-1}$) into realized values of $\hat{x}_t$.
 - ▶ The Analog of b in Bray-Savin is the largest real part of the eigenvalue of B .
- The system is completed by Kalman-updating equations that map $\hat{x}_t$ into $\hat{\mu}_t$.
- Asymptotically $\hat{\mu}_t$ behaves like κt^{b-1} for $\kappa > 0$.

Findings

Table: Eigenvalues of B

	Eigenvalue 1	Eigenvalue 2	T_1	T
ZLB	0.92	-0.48	920,482	944,710
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 on previous slide. The scalar, b , 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.



Findings

Table: Eigenvalues of B

	Eigenvalue 1	Eigenvalue 2	T_1	T
ZLB	0.92	-0.48	920,482	944,710
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 on previous slide. The scalar, b , 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.

- Message: (stable) ZLB has maximal eigenvalue 0.92

Findings

Table: Eigenvalues of B

	Eigenvalue 1	Eigenvalue 2	T_1	T
ZLB	0.92	-0.48	920,482	944,710
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 on previous slide. The scalar, b , 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.

- Message: (stable) ZLB has maximal eigenvalue 0.92
- Asymptotic formulas implies $T_1 = 920,482$ periods to close 2/3 of gap, nonlinear simulations imply $T = 944,710$ periods to converge.

Findings

Table: Eigenvalues of B

	Eigenvalue 1	Eigenvalue 2	T_1	T
ZLB	0.92	-0.48	920,482	944,710
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 on previous slide. The scalar, b , 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.

- Message: (stable) ZLB has maximal eigenvalue 0.92
- Asymptotic formulas implies $T_1 = 920,482$ periods to close 2/3 of gap, nonlinear simulations imply $T = 944,710$ periods to converge.
- When ZLB is ignored, convergence is very fast, and asymptotic formulas reliable far away from REE.

Conclusion

- REE can be a very bad approximation if people in fact are learning.

Conclusion

- REE can be a very bad approximation if people in fact are learning.
- This is particularly true in the ZLB.

Conclusion

- REE can be a very bad approximation if people in fact are learning.
- This is particularly true in the ZLB.
 - ▶ Absence of Taylor principle opens a positive feedback loop from expected π_{t+1} to actual π_{t+1} (b big).

Conclusion

- REE can be a very bad approximation if people in fact are learning.
- This is particularly true in the ZLB.
 - ▶ Absence of Taylor principle opens a positive feedback loop from expected π_{t+1} to actual π_{t+1} (b big).
 - ▶ By learning principle, convergence to REE is slow.

Conclusion

- REE can be a very bad approximation if people in fact are learning.
- This is particularly true in the ZLB.
 - ▶ Absence of Taylor principle opens a positive feedback loop from expected π_{t+1} to actual π_{t+1} (b big).
 - ▶ By learning principle, convergence to REE is slow.
- REE analysis of government spending in ZLB can be very misleading.
 - ▶ Could induce government to rely excessively on fiscal policy.

Conclusion

- REE can be a very bad approximation if people in fact are learning.
- This is particularly true in the ZLB.
 - ▶ Absence of Taylor principle opens a positive feedback loop from expected π_{t+1} to actual π_{t+1} (b big).
 - ▶ By learning principle, convergence to REE is slow.
- REE analysis of government spending in ZLB can be very misleading.
 - ▶ Could induce government to rely excessively on fiscal policy.
- We generalize previous results for asymptotic rates of convergence to vector case and identify analog of Bray and Savin's b .
- Analysis confirms the wisdom of exploring the implication of replacing REE by alternative micro-founded mechanisms by which people form beliefs.

Appendix Materials

Period Price and Profit Functions

- Households (and firms) observe $x = \begin{bmatrix} C, & \pi \end{bmatrix}$
 - ▶ from x (as well as $r, G(r)$) they are able to deduce the variables needed to define their current-period budget constraint.

Period Price and Profit Functions

- Households (and firms) observe $x = \begin{bmatrix} C, & \pi \end{bmatrix}$
 - ▶ from x (as well as $r, G(r)$) they are able to deduce the variables needed to define their current-period budget constraint.
- GDP (Y), aggregate employment (N), real wage (w), marginal firm cost (s), profits, taxes net of profits (T):

$$N = Y = (C + G(r)) \left(1 + \frac{\phi}{2} (\pi - 1)^2 \right)$$

$$w = \chi NC, \quad s = (1 - \nu) w, \quad R = \max \{1, 1 + r^h + \alpha (\pi - 1)\}.$$

Period Price and Profit Functions

- Households (and firms) observe $x = \begin{bmatrix} C, & \pi \end{bmatrix}$
 - ▶ from x (as well as $r, G(r)$) they are able to deduce the variables needed to define their current-period budget constraint.
- GDP (Y), aggregate employment (N), real wage (w), marginal firm cost (s), profits, taxes net of profits (T):

$$N = Y = (C + G(r)) \left(1 + \frac{\phi}{2} (\pi - 1)^2 \right)$$

$$w = \chi NC, \quad s = (1 - \nu) w, \quad R = \max \{1, 1 + r^h + \alpha (\pi - 1)\}.$$

We assume the government issues no debt and finances its expenditures, with lump sum taxes:

$$G(r) + \nu w N,$$

where $\nu w N$ represents the subsidy paid to intermediate good firms.

Period Price and Profit Functions, cnt'd

- Finally, profits net of taxes implied by x and r are:

$$T = \overbrace{(1-s)Y - \frac{\phi}{2}(\pi-1)^2(C+G(r))}^{\text{profits for intermediate good producers}} - \overbrace{(G(r) + \nu wY)}^{\text{lump sum taxes}}.$$



Period Price and Profit Functions, cnt'd

- Finally, profits net of taxes implied by x and r are:

$$T = \overbrace{(1-s)Y - \frac{\phi}{2}(\pi-1)^2(C+G(r))}^{\text{profits for intermediate good producers}} - \overbrace{(G(r) + \nu wY)}^{\text{lump sum taxes}}.$$

- Note: none of these mappings use bond market clearing or the household's intertemporal Euler equation. [▶ Go Back](#)

Cobweb Model

- Model of competitive market and a time lag in production.
 - ▶ John Muth, 'Rational Expectations and the Theory of Price Movements', ECMA, July 1961.
 - ▶ Coase and Fowler, 'Bacon Production and the Pig-Cycle in Great Britain', Economica, May, 1935.
- Demand:

$$d_t = m_I - m_p p_t + v_{1t}$$

- Supply decided in period t before v_{1t} is observed:

$$s_t = r_I + r_p \mathbb{E}_{t-1} p_t + v_{2t}$$

- Equilibrium, $d_t = s_t$:

$$\underbrace{p_t}_{x_t} = \frac{\overbrace{m_I - r_I}^a}{m_p} - \frac{\overbrace{r_p}^{+b}}{m_p} \mathbb{E}_{t-1} p_t + \frac{\overbrace{v_{1t} - v_{2t}}^{\varepsilon_t}}{m_p}$$

Lucas Model

- Aggregate output:

$$q_t = \bar{q} + \pi (p_t - \mathbb{E}_{t-1} p_t) + \zeta_t$$

- Velocity equation:

$$m_t + v_t = p_t + q_t$$

- Monetary policy:

$$m_t = \bar{m} + u_t.$$

- Substitute second two equations into first, to obtain equilibrium condition:

$$\overbrace{p_t}^{x_t} = \frac{\overbrace{\bar{m} - \bar{q}}^a}{1 + \pi} + \frac{\overbrace{\pi}^b}{1 + \pi} \mathbb{E}_{t-1} p_t + \frac{\overbrace{u_t + v_t - \zeta_t}^{\varepsilon_t}}{1 + \pi}$$

Rational Expectations Equilibrium

- Reduced form model:

$$x_t = a + b\mathbb{E}_{t-1}x_t + \varepsilon_t, \varepsilon_t \sim E\varepsilon_t = 0, E\varepsilon_t^2, E\varepsilon_t\varepsilon_{t-j} = 0, j \neq 0.$$

- In rational expectations equilibrium, $\mathbb{E}_{t-1}x_t = E_{t-1}x_t$, so

$$x_t = \frac{a}{1-b} + \varepsilon_t$$

-

Rational Expectations Equilibrium

- Reduced form model:

$$x_t = a + b\mathbb{E}_{t-1}x_t + \varepsilon_t, \varepsilon_t \sim E\varepsilon_t = 0, E\varepsilon_t^2, E\varepsilon_t\varepsilon_{t-j} = 0, j \neq 0.$$

- In rational expectations equilibrium, $\mathbb{E}_{t-1}x_t = E_{t-1}x_t$, so

$$x_t = \frac{a}{1-b} + \varepsilon_t$$

- To verify this, note:

$$\begin{aligned} x_t &= a + bE_{t-1}x_t + \varepsilon_t \stackrel{REE}{=} a + b \overbrace{\frac{a}{1-b}}^{E_{t-1}x_t} + \varepsilon_t \\ &= \frac{a}{1-b} + \varepsilon_t. \end{aligned}$$

Constant-gain learning

- Assume people update their view of μ_{t-1} by constant-gain learning:

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

for $0 < \gamma < 1$.

- Now

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

where $\gamma_b = (1-b)\gamma$,

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

Learning principle again

- Again calculate how long it takes to close $2/3$ of the initial gap, i.e., calculate, T , the value of t such that $z_T \simeq 1/3$.
- Suppose $\gamma = 0.5$ and $b = 0, 0.5, 0.75, 0.85, .95$.

b	0	0.5	0.75	0.85	0.95
T					

Learning principle again

- Again calculate how long it takes to close $2/3$ of the initial gap, i.e., calculate, T , the value of t such that $z_T \simeq 1/3$.
- Suppose $\gamma = 0.5$ and $b = 0, 0.5, 0.75, 0.85, .95$.

b	0	0.5	0.75	0.85	0.95
T	1.6				

Learning principle again

- Again calculate how long it takes to close $2/3$ of the initial gap, i.e., calculate, T , the value of t such that $z_T \simeq 1/3$.
- Suppose $\gamma = 0.5$ and $b = 0, 0.5, 0.75, 0.85, .95$.

b	0	0.5	0.75	0.85	0.95
T	1.6	3.8			

Learning principle again

- Again calculate how long it takes to close $2/3$ of the initial gap, i.e., calculate, T , the value of t such that $z_T \simeq 1/3$.
- Suppose $\gamma = 0.5$ and $b = 0, 0.5, 0.75, 0.85, .95$.

b	0	0.5	0.75	0.85	0.95
T	1.6	3.8	8.23		

Learning principle again

- Again calculate how long it takes to close $2/3$ of the initial gap, i.e., calculate, T , the value of t such that $z_T \simeq 1/3$.
- Suppose $\gamma = 0.5$ and $b = 0, 0.5, 0.75, 0.85, .95$.

b	0	0.5	0.75	0.85	0.95
T	1.6	3.8	8.23	14.1	

Learning principle again

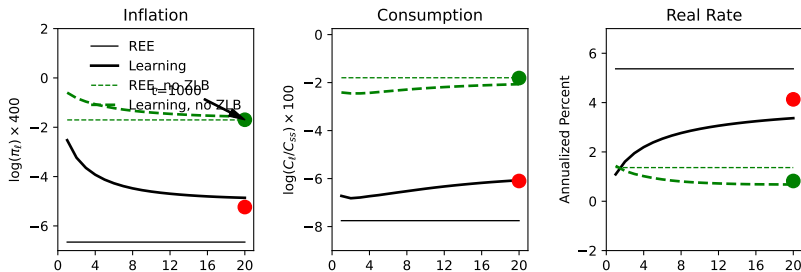
- Again calculate how long it takes to close $2/3$ of the initial gap, i.e., calculate, T , the value of t such that $z_T \simeq 1/3$.
- Suppose $\gamma = 0.5$ and $b = 0, 0.5, 0.75, 0.85, .95$.

b	0	0.5	0.75	0.85	0.95
T	1.6	3.8	8.23	14.1	43.7

- Note: speed of convergence is quicker for 'small' values of b than under Bayesian learning.
- But again speed of convergence increases nonlinearly with b . [▶ Go Back](#)

Dropping ZLB in Experiment #1

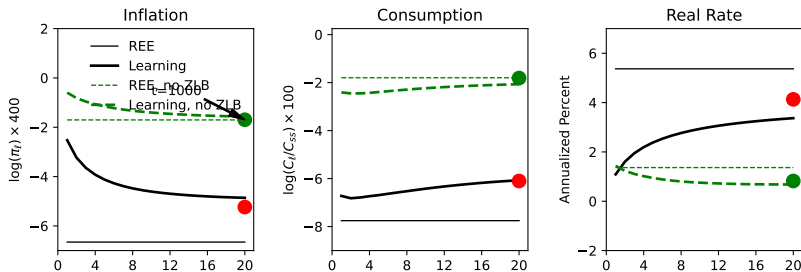
Figure: Effect of Ignoring ZLB



- Ignoring ZLB leads to smaller eventual drop in GDP, as Taylor principle enters picture.

Dropping ZLB in Experiment #1

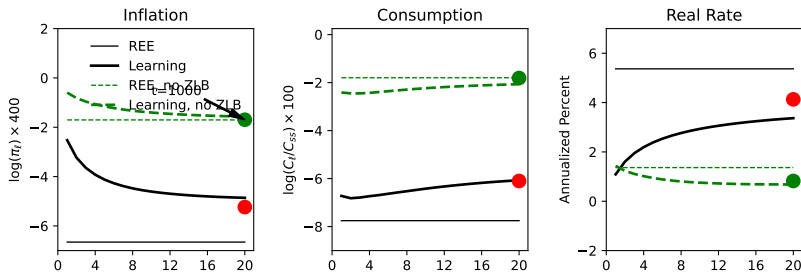
Figure: Effect of Ignoring ZLB



- Ignoring ZLB leads to smaller eventual drop in GDP, as Taylor principle enters picture.
- Seems to move economy in direction of 'low- b economy'.

Dropping ZLB in Experiment #1

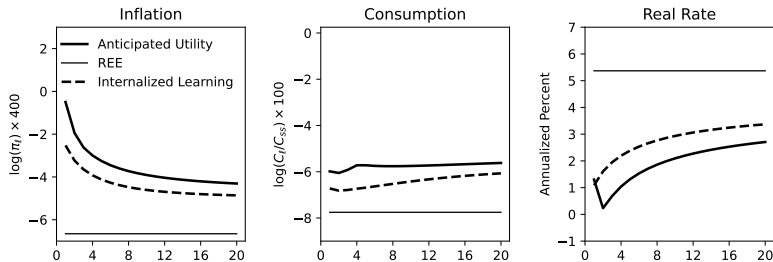
Figure: Effect of Ignoring ZLB



- Ignoring ZLB leads to smaller eventual drop in GDP, as Taylor principle enters picture.
- Seems to move economy in direction of 'low- b economy'.
- Still studying this result. [▶ Go Back](#)

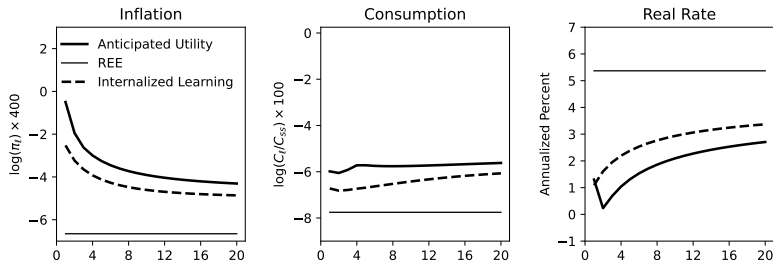
Role of Internalized Learning in Experiment #1

- r drops and G remains unchanged.



Role of Internalized Learning in Experiment #1

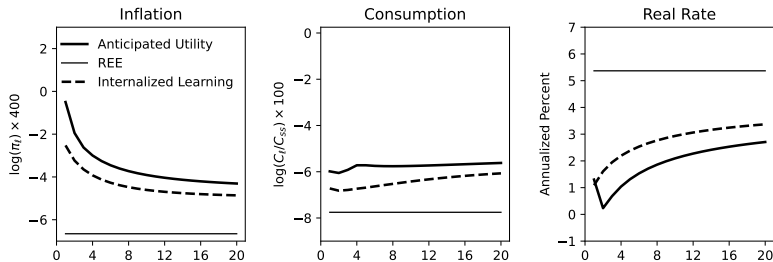
- r drops and G remains unchanged.



- Qualitatively, results not dependent on whether we use anticipated utility or internalized learning.
 - ▶ Very slow convergence in each case.

Role of Internalized Learning in Experiment #1

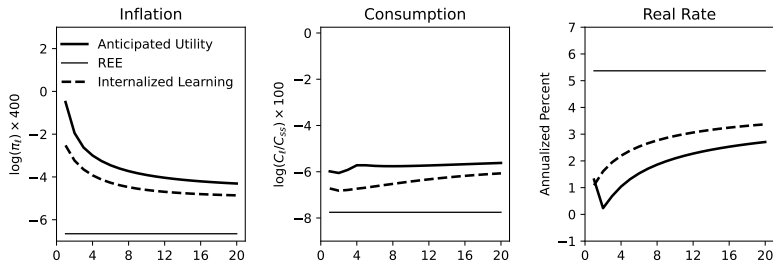
- r drops and G remains unchanged.



- Qualitatively, results not dependent on whether we use anticipated utility or internalized learning.
 - ▶ Very slow convergence in each case.
- Reasons for the difference:
 - ▶ Prices actually fall in period 0, and under internalized learning people's beliefs respond to that.

Role of Internalized Learning in Experiment #1

- r drops and G remains unchanged.



- Qualitatively, results not dependent on whether we use anticipated utility or internalized learning.
 - ▶ Very slow convergence in each case.
- Reasons for the difference:
 - ▶ Prices actually fall in period 0, and under internalized learning people's beliefs respond to that.
 - ▶ Precautionary motive stronger under greater uncertainty of internalized learning.

[▶ Go Back](#)

Multiple REE in ZLB

- Scenario:
 - ▶ Economy was in zero inflation steady state up to period 0
 - ▶ Unexpectedly, discount rate shock happens and everyone correctly believes that economy goes back to zero inflation steady state with constant probability.

Multiple REE in ZLB

- Scenario:

- ▶ Economy was in zero inflation steady state up to period 0
- ▶ Unexpectedly, discount rate shock happens and everyone correctly believes that economy goes back to zero inflation steady state with constant probability.
- ▶ Well known: there are two stationary rational expectations ZLB equilibria.
 - ★ In our model, can characterize a ZLB equilibrium as a zero of a function of inflation alone, $f(\pi_\ell) = 0$.
 - ★ This function has an 'inverse U', Laffer curve shape.

Multiple REE in ZLB

- Scenario:

- ▶ Economy was in zero inflation steady state up to period 0
- ▶ Unexpectedly, discount rate shock happens and everyone correctly believes that economy goes back to zero inflation steady state with constant probability.
- ▶ Well known: there are two stationary rational expectations ZLB equilibria.
 - ★ In our model, can characterize a ZLB equilibrium as a zero of a function of inflation alone, $f(\pi_\ell) = 0$.
 - ★ This function has an 'inverse U', Laffer curve shape.

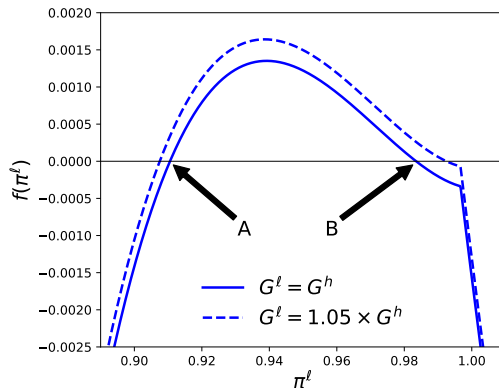
- Parameter values

$$p = 0.80, r_\ell = -0.0015 (-0.6APR), G_{ss} = 0.20, r_{ss} = 0.005 (2.0APR),$$

$$Y_{ss} = N_{ss} = 1, \varepsilon = 7, \phi = 110, \chi = 1.25, \alpha = 1.5$$

REE Equilibria in ZLB

- Two ZLB equilibria
 - ▶ Bad-ZLB (A) equilibrium: substantial deflation, very high real rate, very low consumption.
 - ▶ Good-ZLB (B) equilibrium: more modest deflation, reduced consumption and high in real rate.



Does Learning Select One of the Two Equilibria?

Does Learning Select One of the Two Equilibria?

- Bad-ZLB equilibrium is locally unstable under learning.
 - ▶ When prior means are centered (priors on variance positive) on Bad-ZLB, you go to Good-ZLB.

Does Learning Select One of the Two Equilibria?

- Bad-ZLB equilibrium is locally unstable under learning.
 - ▶ When prior means are centered (priors on variance positive) on Bad-ZLB, you go to Good-ZLB.
- Good-ZLB equilibrium is 'globally' stable under learning.
 - ▶ When prior mean of x is centered on steady state, on Good-ZLB or on Bad-ZLB:

Does Learning Select One of the Two Equilibria?

- Bad-ZLB equilibrium is locally unstable under learning.
 - ▶ When prior means are centered (priors on variance positive) on Bad-ZLB, you go to Good-ZLB.
- Good-ZLB equilibrium is 'globally' stable under learning.
 - ▶ When prior mean of x is centered on steady state, on Good-ZLB or on Bad-ZLB: converge to Good-ZLB. [▶ Go Back](#)