

Measuring Uncertainty about Long-Run Predictions*

Ulrich K. Müller and Mark W. Watson

Princeton University

Department of Economics

Princeton, NJ, 08544

February 2013 (Revised September 2014)

Abstract

Long-run forecasts of economic variables play an important role in policy, planning, and portfolio decisions. We consider forecasts of the long-horizon average of a scalar variable, typically the growth rate of an economic variable. The main contribution is the construction of prediction sets with asymptotic coverage over a wide range of data generating processes, allowing for stochastically trending mean growth, slow mean reversion and other types of long-run dependencies. We illustrate the method by computing prediction sets for 10 to 75 year average growth rates of U.S. real per-capita GDP, consumption, productivity, price level, stock prices and population.

JEL classification: C22, C53, E17

Keywords: prediction interval, low frequency, spectral analysis, least favorable distribution

*We thank Frank Diebold, Graham Elliott, Bruce Hansen, James Stock, Jonathan Wright, three referees, and the Editor, Stéphane Bonhomme, for useful comments and advice. Support was provided by the National Science Foundation through grants SES-0751056 and SES-1226464.

1 Introduction

This paper is concerned with quantifying the uncertainty in long-run predictions of economic variables. Long-run forecasts and the uncertainty surrounding them play an important role in policy, planning, and portfolio decisions. For example, in the United States, an ongoing task of the Congressional Budget Office (CBO) is to forecast productivity and real GDP growth over a 75-year horizon to help gauge the solvency of the Social Security Trustfund. Uncertainty surrounding these forecasts is then translated into the probability of trust fund insolvency.¹ Inflation “Caps” and “Floors” are option-like derivatives with payoffs tied to the average value of price inflation over the next decade; their risk-neutral prices are determined by the probability that the long-run average of future values of inflation falls above or below a pre-specified threshold.² And, there is a large literature in finance discussing optimal portfolio allocations for long-run investors and how these portfolios depend on uncertainty in long-run returns.³

Let x_t denote a time series, such as the inflation rate, the growth rate of real GDP or the return on a portfolio of stocks. Sample data on x_t are available for $t = 1, \dots, T$, say 1947-2014. Let $\bar{x}_{T+1:T+h} = h^{-1} \sum_{t=1}^h x_{T+t}$ denote the average value of the series between time periods $T + 1$ through $T + h$, say the 32-year horizon 2014-2046. We are interested in the date T uncertainty about the value of $\bar{x}_{T+1:T+h}$, as characterized by prediction sets that contain $\bar{x}_{T+1:T+h}$ with a pre-specified probability (such as 90%). This is a long-horizon problem, since the horizon h is large relative to the number of available observations T (in the example, $r = h/T \approx 0.5$).

We structure the problem so that the coverage probability can be calculated using asymptotic approximations based on a central limit theorem. In particular we suppose that both T and h are large, and construct the prediction sets as a function of a relatively small number of weighted averages of the sample values of x_t . We apply a central limit theorem to the variable of interest ($\bar{x}_{T+1:T+h}$) and the predictors, and study an asymptotic version of the prediction problem based on the multivariate normal distribution. Were all the parameters of this normal distribution known (or consistently estimable), the prediction problem would be a straightforward application of optimal prediction in the multivariate normal model.

The problem is complicated by unknown parameters that characterize the stochastic process x_t and hence also the covariance matrix of the normal distribution in the large-sample

¹See Congressional Budget Office (2005).

²See Fleckenstein, Longstaff, and Lustig (2013), Hilsher, Raviv and Reis (2014), and Kitsul and Wright (2012) who use market prices on various inflation-related derivatives to estimate market-based predictive distributions of inflation.

³See, for example, Campbell and Viceira (1999), Pastor and Stambaugh (2012), and Siegel (2007).

problem. We assume that the first differences $\Delta x_t = x_t - x_{t-1}$ are covariance stationary. (Recall that x_t is a series like the growth rate of real GDP, inflation, or asset returns, so this does not rule out stochastic trends in these *growth rates*.) Since we are interested in a long-run prediction ($\bar{x}_{T+1:T+h}$, for h large relative to T) the crucial characteristic of x_t is its (pseudo-) spectrum near frequency zero. The relative paucity of sample information about these low-frequency properties precludes a nonparametric approach. We therefore proceed by constructing a flexible parametric model for the shape of the spectrum near frequency zero that nests the fractional, local-to-unity and local-level forms of long run persistence. The uncertainty about the parameter θ of this model in turn becomes an important component of the uncertainty about $\bar{x}_{T+1:T+h}$.

We use both Bayes and frequentist methods to incorporate this uncertainty in our prediction sets. The Bayes procedure is straightforward: given a prior for the parameter θ , and the Gaussianity of the limiting problem, the predictive density for $\bar{x}_{T+1:T+h}$ follows from Bayes rule, so that prediction sets are readily computed. While Bayes sets have many desirable properties, they have the potentially undesirable property of controlling coverage (that is, the probability that the set includes the future value of $\bar{x}_{T+1:T+h}$) only *on average* for values of θ drawn from the prior. Thus in general, coverage will fall short of this average value for some values of θ , and the specifics of this undercoverage will depend on the prior used. To address this limitation we robustify the Bayes prediction sets by enlarging them so that they have frequentist properties: the resulting sets provide (possibly conservative) coverage for all values of θ . Using ideas borrowed from Müller and Norets (2012), we do this in a way that minimizes the sets' average expected length.

In economics, arguably the most well-known predictive densities and corresponding prediction sets are the “Rivers of Blood” shown in the Bank of England’s *Inflation Report*. These are judgmental prediction sets for inflation that are computed over a four year horizon by the members of the Bank’s Monetary Policy Committee. In contrast, we are interested in prediction sets computed from probability models over long horizons, and the literature on this topic is relatively sparse. Most of the existing literature on long-horizon forecasting stresses the difficulty of constructing good long-term forecasts under uncertainty about the long-run properties of the process. Granger and Jeon (2007) provide a mostly verbal account. Elliott (2006) compares alternative approaches to point forecasts and compares their mean squared errors. Kemp (1999), Phillips (1998) and Stock (1996, 1997) show that standard formulas for forecast uncertainty break down in the long-horizon local-to-unity model, but they do not provide constructive alternatives. In the related problem of estimating long-run impulse responses Pesavento and Rossi (2006) construct confidence sets that account for uncertainty about the local-to-unity parameter. Chapter 8.7 in Beran (1994) discusses

forecasting of fractionally integrated series, and Doornik and Ooms (2004) use an ARFIMA model to generate long-run uncertainty bands for future inflation, but without accounting for parameter estimation uncertainty. Two strands of literature study long-run forecast uncertainty for time series that we analyze by constructing series-specific Bayesian models. Pastor and Stambaugh (2012) compute predictive variances of long-run forecasts of stock returns that account for parameter uncertainty. Lee (2011) and Raftery, Li, Sevcikova, Gerland, and Heilig (2012) study long-run forecasts of population and fertility rates.

The outline of this paper is as follows. Section 2 formalizes the long-horizon prediction problem and discusses the low-frequency summaries of the sample data used in the analysis. This section also introduces two running examples: forecasting the average growth rate of real per-capita GDP and the average level of price inflation in the U.S. over the next 25 years. Section 3 discusses and develops the requisite statistical tools for constructing the long-horizon prediction sets. Two sets of tools are needed. The first is a central limit theorem and associated covariance matrix that yields a large-sample Gaussian version of the prediction problem. The second are methods for constructing Bayes and frequentist prediction sets for this limiting problem. The Bayes procedures are standard; the frequentist procedures are not, and are developed in Section 3.3. Section 4 takes up the important practical problems of parameterizing the covariance matrix in the limiting problem (which involves parameterizing the spectrum of x_t near frequency 0), choosing a prior for the Bayes prediction sets and a related weighting function for the frequentist sets (to obtain a scalar criterion for comparing the efficiency of sets), and choosing the number of low-frequency averages of the sample data to use (which involves a classic trade-off between efficiency and robustness). Taken together, Sections 2-4 develop methods for constructing prediction sets with well-defined large-sample optimality properties; these methods are illustrated using the GDP and inflation running examples throughout these sections. Section 5 uses simulations and pseudo-out-of-sample experiments to evaluate the performance of these sets in small samples. One focus of this analysis is the effect of level or volatility "breaks" on the prediction sets. Following this extensive background, Section 6 applies these methods to construct prediction sets spanning up to 75 years for the eight U.S. economic time series, the running examples of real GDP and inflation (as measured by the CPI), but also the rates of growth of per-capita consumption expenditures, productivity (total factor and labor productivity), population, stock prices, and alternative measures of price inflation. Section 7 concludes.

2 The Prediction Problem

Let x_t be the economic variable of interest which is observed for $t = 1, \dots, T$. The objective is to construct a prediction set, denoted by A , of the average value of x_t from periods $T + 1$ to $T + h$,

$$\bar{x}_{T+1:T+h} = h^{-1} \sum_{t=1}^h x_{T+t} \quad (1)$$

with the property that $P(\bar{x}_{T+1:T+h} \in A) = 1 - \alpha$, where α is a pre-specified constant. The prediction set A is constructed using the sample data for x_t , so that $A = A(\{x_t\}_{t=1}^T)$.⁴ We restrict A in two ways. First, we allow A to depend on the sample data only through a small number low-frequency weighted averages of the sample data, and second, we restrict A to be scale and location equivariant. We discuss each of these restrictions in turn.

Cosine transformations of the sample data. Because h is large, the prediction sets involve long-run uncertainty about x_t . It is therefore useful to transform the sample data into weighted averages that capture variability at different frequencies – we will be interested in the weighted averages corresponding to low frequencies. Thus, consider the weighted averages $(\bar{x}_{1:T}, X_T)$, with $\bar{x}_{1:T} = T^{-1} \sum_{t=1}^T x_t$, $X_T = (X_T(1), \dots, X_T(T-1))'$, and where $X_T(j)$ is the j th cosine transformation

$$X_T(j) = \int_0^1 \Psi_j(s) x_{\lfloor sT \rfloor + 1} ds = \iota_{jT} T^{-1} \sum_{t=1}^T \Psi_j\left(\frac{t-1/2}{T}\right) x_t \quad (2)$$

with $\Psi_j(s) = \sqrt{2} \cos(j\pi s)$ and $\iota_{jT} = (2T/j\pi) \sin(j\pi/2T) \rightarrow 1$. The cosine transforms have two properties we will exploit. First, they isolate variation in the sample data corresponding to different frequencies: $\bar{x}_{1:T}$ captures 0-frequency variation and $X_T(j)$ captures variation at frequency $j\pi/T$. Second, because the Ψ_j weights add to zero, $X_T(j)$ is invariant to location shifts of the sample, a property we use when we construct equivariant prediction sets.

The $T \times 1$ vector $(\bar{x}_{1:T}, X_T)$ is a nonsingular transformation of the sample data $\{x_t\}_{t=1}^T$, but we will construct prediction sets based on a truncated information set that includes only $\bar{x}_{1:T}$ and the first q cosine transforms, $X_{T,1:q} = (X_T(1), X_T(2), \dots, X_T(q))'$ and where q is much smaller than $T - 1$. Thus, the prediction sets we consider are $A = A(\bar{x}_{1:T}, X_{T,1:q})$, and so rely solely on variability in the data associated with frequencies lower than $q\pi/T$. We

⁴Of course, when x_t is the first difference of another variable y_t , so that $x_t = y_t - y_{t-1}$, then forecasts of y_{T+h} can be constructed from forecasts of $\bar{x}_{T+1:T+h}$ using the identity $y_{T+h} = y_T + h\bar{x}_{T+1:T+h}$. Moreover, prediction sets for $\bar{x}_{T+1:T+h}$ and y_{T+h} are readily converted into prediction sets for monotonic transformation of these variables. For example, a prediction set for the average growth rate of real GDP ($\bar{x}_{T+1:T+h}$) yields a prediction set for the log-level of real GDP (y_{T+h}) or the level of real GDP ($\exp(y_{T+h})$).

compress the sample information into the $q + 1$ variables $(\bar{x}_{1:T}, X_{T,1:q})$ for two reasons. The first is tractability: with a focus on this truncated information set, the analysis involves a small number of variables (the $(q + 2)$ variables $(\bar{x}_{T+1:T+h}, \bar{x}_{1:T}, X_{T,1:q})$), and because each of these variables is a weighted average of $\{x_t\}_{t=1}^{T+h}$, a central limit theorem derived in the next section allows us to study a limiting Gaussian version of the prediction problem that is much simpler than the original finite-sample problem. The second motivation for truncating the information set is robustness: we use the low-frequency information in the sample data ($\bar{x}_{1:T}$ and the first q elements of X_T) to inform us about a low-frequency, long-run average of future data, but we do not use high frequency sample information (the last $T - 1 - q$ elements of X_T). While high frequency information is informative about low-frequency characteristics for some stochastic processes (for example, tightly parameterized ARMA processes), this is generally not the case, and high-frequency sample variation may lead to faulty low-frequency inference. Müller and Watson (2008, 2013) discuss this issue in detail. In Section 4 below we present numerical calculations that quantify the efficiency-robustness trade-off embodied by the choice of q in the long-run prediction problem.

Invariance. In our applications it is natural to restrict attention to prediction sets that are invariant to location and scale, so for example, the results will not depend on whether the data are expressed as growth rates in percentage points at an annual rate or as percent per quarter. Thus, we restrict attention to prediction sets with the property that if $y \in A(\bar{x}_{1:T}, X_{T,1:q})$ then $m + by \in A(m + b\bar{x}_{1:T}, bX_{T,1:q})$ for any constants m and $b \neq 0$ (where the transformation of $X_{T,1:q}$ does not depend on m because, as mentioned above, $X_{T,1:q}$ is location invariant). Invariance allows us to restrict attention to prediction sets that depend on functions of the sample data that are scale and location invariant; in particular we can limit attention to constructing prediction sets for Y_T^s given $X_{T,1:q}^s$, where $Y_T^s = Y_T / \sqrt{X_{T,1:q}' X_{T,1:q}}$ with

$$Y_T = \bar{x}_{T+1:T+h} - \bar{x}_{1:T} \quad (3)$$

and $X_{T,1:q}^s = X_{T,1:q} / \sqrt{X_{T,1:q}' X_{T,1:q}}$.⁵

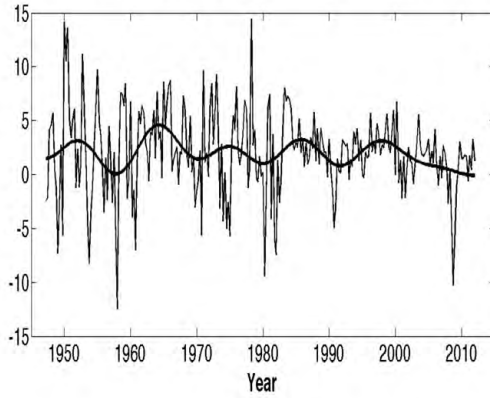
Running examples: Two of the economic time series studied in Section 6 are the growth rate of U.S. real per-capita GDP and the rate of inflation in the U.S. based on the consumer price index. We use these series as running examples to illustrate concepts as they are introduced. Panels (i) in Figure 1 plot the quarterly values of these time series from 1947-2012, along with the low-frequency components of the time series formed as the projection of

⁵Setting $m = -\bar{x}_{1:T} / \sqrt{X_{T,1:q}' X_{T,1:q}}$ and $b = 1 / \sqrt{X_{T,1:q}' X_{T,1:q}}$ implies that for any invariant set A , $y \in A(\bar{x}_{1:T}, X_{T,1:q})$ if and only if $(y - \bar{x}_{1:T}) / \sqrt{X_{T,1:q}' X_{T,1:q}} \in A(0, X_{T,1:q} / \sqrt{X_{T,1:q}' X_{T,1:q}})$, and thus also $\bar{x}_{T+1:T+h} \in A(\bar{x}_{1:T}, X_{T,1:q})$ if and only if $Y_T^s \in A(0, X_{T,1:q}^s)$.

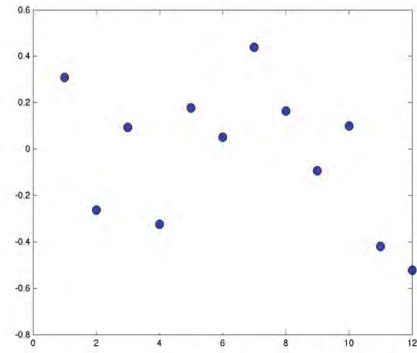
Figure 1: Low-frequency components and cosine transforms

(a) Growth rate of real per-capita GDP

(i) Series and low-frequency component

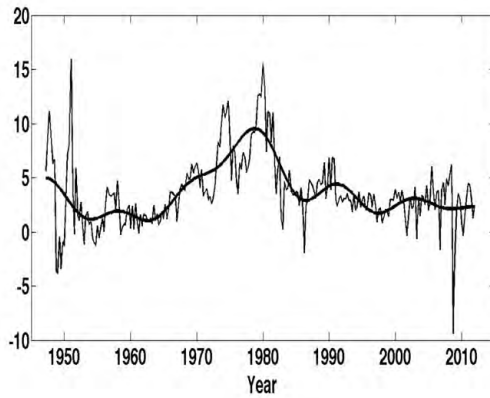


(ii) Cosine transform

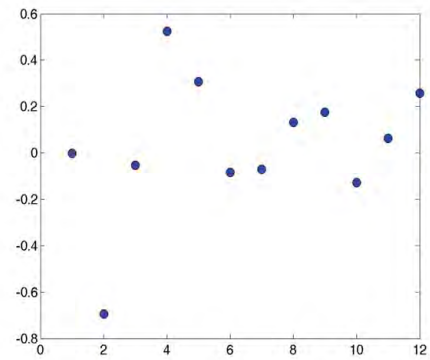


(b) Inflation (CPI)

(i) Series and low-frequency component



(ii) Cosine transform



Notes: The low-frequency components in (i) are the projection of the series onto $\cos[(t-0.5)\pi j/T]$ for $j = 0, \dots, 12$. The cosine transforms shown in (ii) are the standardized values, $X_{T,tj}^s$.

the series onto $\cos[(t - 1/2)\pi j/T]$ for $j = 0, \dots, 12$. (The value $q = 12$ is used in the empirical analysis in Section 6 for reasons discussed below). The coefficients in the projection are the cosine transformations, $X_{T,1:q}$, and their standardized values, $X_{T,1:q}^s$ are plotted in panels (ii). These low-frequency components of the data are the summaries of the sample data we use to construct long-horizon prediction sets. Looking at panels (i), inflation exhibits much more low-frequency variation than GDP growth rates over the sample period; this is manifested in panels (ii) by the relatively larger magnitude of inflation's first few cosine transformations, capturing pronounced low-frequency movements. $\blacktriangle$

3 Statistical Preliminaries

The last section laid out the finite-sample prediction problem. In this section we review and develop the statistical theory that will guide our approach to constructing prediction sets. We divide the section into three subsections. The first provides a central limit theorem that characterizes the large-sample behavior of the weighted averages $(X_{T,1:q}, Y_T)$, and provides a characterization of the limiting covariance matrix based on the properties of the (pseudo-) spectrum for x_t near frequency zero. The second subsection illustrates this framework in the fractional $I(d)$ model and reports prediction sets for known d and Bayes prediction sets using a prior for d . The final subsection discusses the generic problem of robustifying Bayes prediction sets to obtain sets with frequentist coverage uniformly over the parameter space.

3.1 Large-Sample Approximations

3.1.1 Asymptotic Behavior of $(X_{T,1:q}, Y_T)$

To derive the asymptotic behavior for $(X_{T,1:q}, Y_T)$, note that each element can be written as a weighted average of the elements of $\{x_t\}_{t=1}^{T+h}$. Thus, let $g : [0, 1 + r] \mapsto \mathbb{R}$ denote a generic weighting function, where $r = \lim_{T \rightarrow \infty} (h/T) > 0$, and consider the weighted average of $\{x_t\}_{t=1}^{T+h}$

$$\eta_T = T^{1-\alpha} \int_0^{1+r} g(s) x_{\lfloor sT \rfloor + 1} ds$$

where α is a suitably chosen constant.⁶ In our context, the elements of $X_{T,1:q}$ are cosine transformations of the in-sample values of x_t (cf. (2)), so that $g(s) = \sqrt{2} \cos(j\pi s)$ for $0 \leq s \leq 1$ and $g(s) = 0$ for $s > 1$; Y_T defined in (3) is the difference between the out-of-sample and in-sample average values of x_t , so that $g(s) = -1$ for $0 \leq s \leq 1$ and $g(s) = r^{-1}$

⁶For the $I(0)$ model $\alpha = 1/2$, for the $I(1)$ model $\alpha = 3/2$, and so forth.

for $1 < s \leq 1 + r$. These weights sum to zero, so that the (unconditional) expectation of x_t plays no role in the study of η_T .

In the appendix we provide a central limit theorem for η_T under a set of primitive conditions about the stochastic process describing x_t and these weighting functions. We will not list the technical conditions in the text, but rather give a brief overview of the key conditions before stating the limiting result and discussing the form of the limiting covariance matrix. In particular, the analysis is carried out under the assumption that $\Delta x_t = \Delta x_{T,t}$ is a double array process with moving average representation $\Delta x_{T,t} = c_T(L)\varepsilon_t$, where ε_t is a possibly conditionally heteroskedastic martingale difference sequence with more than 2 unconditional moments.⁷ The moving average coefficients in $c_T(L)$ are square summable for each T , so that $\Delta x_{T,t}$ has a spectrum, denoted by $F_T(\lambda)$. The motivation for allowing $c_T(L)$ and F_T to depend on T is to capture many forms of persistence, as stemming from an autoregressive root local-to-unity, $\rho_T = 1 - c/T$, for instance.

Under these and additional technical assumptions, Theorem 1 in the appendix shows that η_T has a limiting normal distribution,⁸ and as an implication

$$T^{1-\alpha} \begin{bmatrix} X_{T,1:q} \\ Y_T \end{bmatrix} \Rightarrow \begin{bmatrix} X \\ Y \end{bmatrix} \sim \mathcal{N}(0, \Sigma) \sim \mathcal{N} \left(0, \begin{pmatrix} \Sigma_{XX} & \Sigma_{XY} \\ \Sigma_{YX} & \Sigma_{YY} \end{pmatrix} \right) \quad (4)$$

with $X = (X_1, \dots, X_q)'$ (we omit the dependence of X on q to ease notation), and Σ a function of the limiting properties of the spectrum F_T near zero (see the next subsection).

The limiting density of the invariants $X_{T,1:q}^s = X_{T,1:q} / \sqrt{X_{T,1:q}' X_{T,1:q}}$ and $Y_T^s = Y_T / \sqrt{X_{T,1:q}' X_{T,1:q}}$ follows directly from (4) and the continuous mapping theorem,

$$\begin{bmatrix} X_{T,1:q}^s \\ Y_T^s \end{bmatrix} \Rightarrow \begin{bmatrix} X^s \\ Y^s \end{bmatrix} = \begin{bmatrix} X / \sqrt{X' X} \\ Y / \sqrt{X' X} \end{bmatrix}. \quad (5)$$

Note that as a consequence of the imposed scale invariance, the convergence (5) does not involve α , and the distribution of (X^s, Y^s) does not depend on the scale of Σ . Explicit expressions for the densities f_{X^s} and $f_{(X^s, Y^s)}$ of X^s and (X^s, Y^s) as a functions of Σ_{XX} and Σ are provided in the appendix.

⁷The restriction that $E\Delta x_t = 0$ rules out a deterministic trend in x_t . This restriction is plausible in our empirical analysis in which x denotes growth rates of real variables like per capita GDP, inflation rates, and asset returns.

⁸As in any central limit theorem, the conditions underlying Theorem 1 imply that no single shock has a substantial impact on the overall variability of η_T . This assumption might be violated by rare but catastrophic events stressed in the work of Rietz (1988) and Barro (2006), for example. Note, however, that such events would need to substantially impact the *average* growth rate over a long horizon to invalidate a normal approximation.

With Σ known, it is straightforward to compute prediction sets of Y^s given $X^s = x^s$: A calculation shows that the distribution of Y^s conditional on $X^s = x^s$ satisfies (see the appendix)

$$\frac{Y^s - \Sigma_{YX}\Sigma_{XX}^{-1}x^s}{\sqrt{\Sigma_{YY} - \Sigma_{YX}\Sigma_{XX}^{-1}\Sigma_{XY}}\sqrt{x^{s'}\Sigma_{XX}^{-1}x^s/q}} \sim \text{Student-}t^q \quad (6)$$

so that prediction sets for Y^s of a given level $1 - \alpha$ are readily computed using Student- t quantiles. These sets in turn imply asymptotically justified prediction sets for Y_T^s via (5), and thus also for $\bar{x}_{T+1:T+h}$ via the definition of $(X_{T,1;q}^s, Y_T^s)$ and (3).

In particular, when $x_{T,t}$ is $I(0)$ with long-run variance σ^2 , it turns out that $\Sigma_{YX} = 0$, $\Sigma_{XX} = \sigma^2 I_q$, $\Sigma_{YY} = (1 + r^{-1})\sigma^2$, and the $1 - \alpha$ prediction set for Y is given by the interval with endpoints $\pm t_{(1-\alpha/2)}^q \sqrt{q(1 + r^{-1})X'X}$, where $t_{(1-\alpha/2)}^q$ is the $(1 - \alpha/2)$ quantile of the t -distribution of q degrees of freedom. The prediction set for $\bar{x}_{T+1:T+h}$ is therefore $\bar{x}_{1:T} \pm t_{(1-\alpha/2)}^q (1 + r^{-1})^{1/2} T^{-1/2} s_{LR}$, where $s_{LR}^2 = (T/q) X'_{T,1;q} X_{T,1;q}$.

3.1.2 The Local-to-Zero Spectrum

The appendix derives Σ , the covariance matrix in (4), and shows that it depends exclusively on the low-frequency properties of x_t (we omit the dependence on T to ease notation). Specifically, with $R_T(\lambda) = F_T(\lambda)/|1 - e^{-i\lambda}|^2$ the (pseudo-) spectrum of x_t , let

$$S(\omega) = \lim_{T \rightarrow \infty} T^{1-2\alpha} R_T(\omega/T)$$

denote the "local-to-zero" spectrum of x_t . The covariance matrix Σ depends on $S(\omega)$ and the weighting functions $g(s)$ used for the sample averages. The details are provided in the appendix, but a calculation for the special case in which x_t is stationary (so that R_T is its spectrum) demonstrates the result.

Thus, suppose that x_t is stationary. The jk 'th element of Σ is the limiting covariance of the two weighted averages $\eta_{j,T}$ and $\eta_{k,T}$, where

$$\eta_{l,T} = T^{1-\alpha} \int_0^{1+r} g_l(s) x_{[sT]+1} ds = T^{-\alpha} \sum_{t=1}^{\lfloor (1+r)T \rfloor} \tilde{g}_{l,t} x_t$$

with $\tilde{g}_{l,t} = T \int_{(t-1)/T}^{t/T} g_l(s) ds$. Recalling that the j th autocovariance of x_t is given by $\int_{-\pi}^{\pi} R_T(\lambda) e^{-i\lambda j} d\lambda$, where $i = \sqrt{-1}$, we obtain

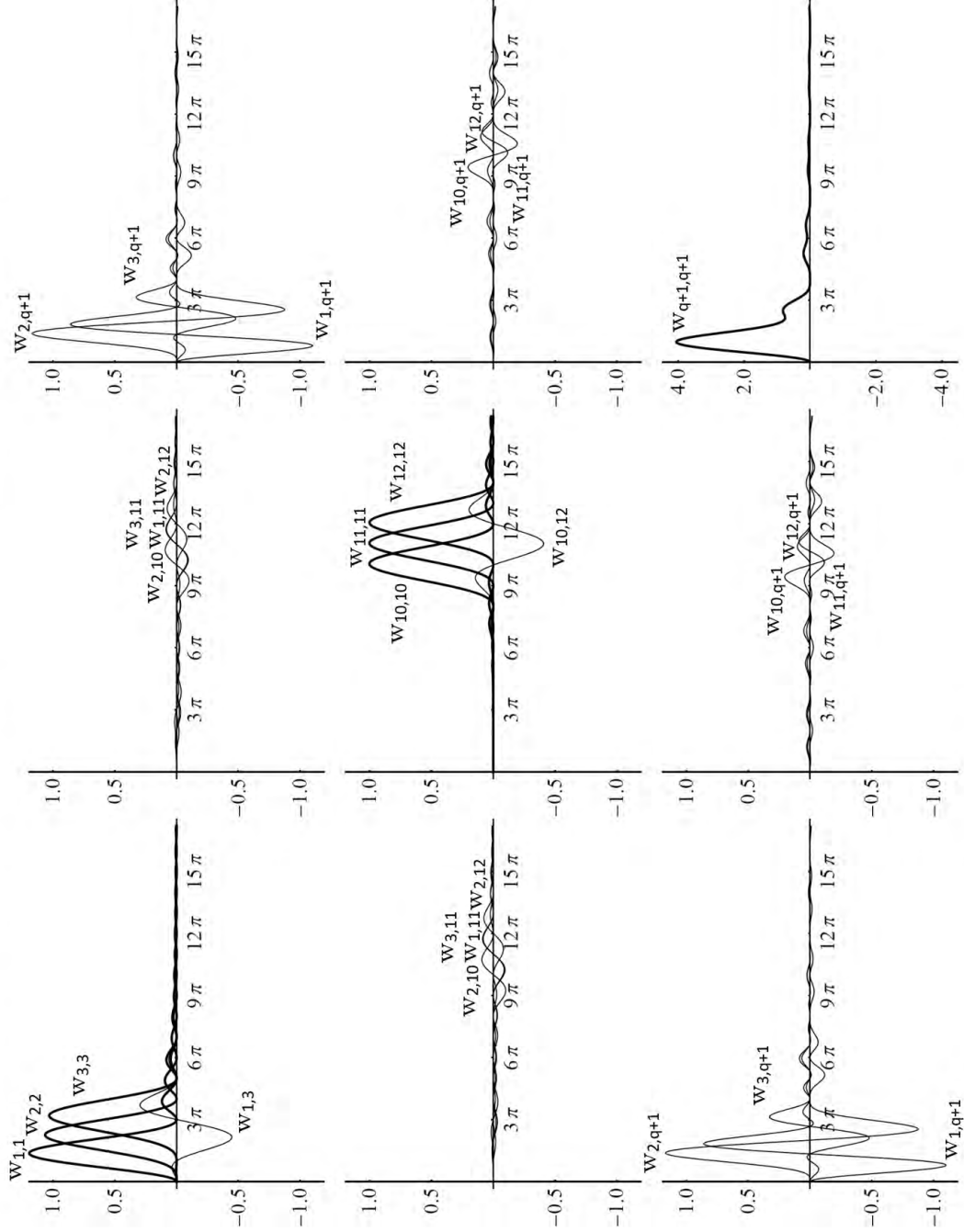
$$E(\eta_{j,T} \eta_{k,T}) = T^{-2\alpha} \sum_{s,t=1}^{\lfloor (1+r)T \rfloor} \left(\int_{-\pi}^{\pi} e^{-i\lambda(t-s)} R_T(\lambda) d\lambda \right) \tilde{g}_{j,t} \tilde{g}_{k,s}$$

$$\begin{aligned}
&= T^{-2\alpha} \int_{-\pi}^{\pi} R_T(\lambda) \left(\sum_{s=1}^{\lfloor (1+r)T \rfloor} \tilde{g}_{k,s} e^{i\lambda s} \right) \left(\sum_{t=1}^{\lfloor (1+r)T \rfloor} \tilde{g}_{j,t} e^{-i\lambda t} \right) d\lambda \\
&= T^{1-2\alpha} \int_{-T\pi}^{T\pi} R_T(\omega/T) \left(T^{-1} \sum_{s=1}^{\lfloor (1+r)T \rfloor} \tilde{g}_{k,s} e^{i\omega(s/T)} \right) \left(T^{-1} \sum_{t=1}^{\lfloor (1+r)T \rfloor} \tilde{g}_{j,t} e^{-i\omega(t/T)} \right) d\omega \\
&\rightarrow \int_{-\infty}^{\infty} S(\omega) \left(\int_0^{1+r} g_k(s) e^{i\omega s} ds \right) \left(\int_0^{1+r} g_j(s) e^{-i\omega s} ds \right) d\omega = \Sigma_{j,k}. \tag{7}
\end{aligned}$$

The limit (7) shows that the elements of the covariance matrix of $(X, Y)'$ are simply weighted averages of the local-to-zero spectrum S .

Several features of Σ follow from (7). For example, since S is an even function and g_j and g_k are real valued, (7) can be rewritten as $E(\eta_{j,T} \eta_{k,T}) \rightarrow 2 \int_0^\infty S(\omega) w_{jk}(\omega) d\omega$, where $w_{jk}(\omega) = \text{Re}[\left(\int_0^{1+r} g_j(s) e^{-i\omega s} ds\right) \left(\int_0^{1+r} g_k(s) e^{i\omega s} ds\right)]$. With $g_j = \sqrt{2} \cos(\pi j s)$, a calculation shows that $w_{jk}(\omega) = 0$ for $1 \leq j, k \leq q$ and $j + k$ odd, so that $E(X_j X_k) = 0$ for all odd $j + k$, independent of the local-to-zero spectrum S . Figure 2 plots $w_{jk}(\cdot)$ for some selected values of j and k and $r = 1/2$. The figure displays the weights $w_{j,k}$ corresponding to the covariance matrix of the vector $(X_1, X_2, X_3, X_{10}, X_{11}, X_{12}, Y)'$, organized into a symmetric matrix of 9 panels. The first panel plots the weights for the covariance matrix of $(X_1, X_2, X_3)'$, where the weights for the variances are shown in bold and the weights for the covariances are shown as thin curves. The second panel plots the weights associated with covariances between $(X_1, X_2, X_3)'$ and $(X_{10}, X_{11}, X_{12})'$, and so forth. The covariance weights w_{jk} , $j \neq k$, integrate to zero, which implies that for a flat local-to-zero spectrum S (corresponding to an $I(0)$ model), Σ is diagonal. As can be seen from the first and second panel on the diagonal of Figure 1, the variance of the predictor X_j is mostly determined by the values of S in the interval $\pi j \pm 2\pi$. Further, as long as S is somewhat smooth, the correlation between X_j and X_k is very close to zero for $|j - k|$ large. The final panel in the figure shows the weight associated with the variance of Y , the variable being predicted. Evidently the unconditional variance of Y is mostly determined by the shape of S on the interval $\omega \in [0, 4\pi]$, and its correlation with X_j is therefore small for j large even for smooth but non-constant S (a calculation shows that $\max_\omega |w_{q+1,j}(\omega)|$ decays at the rate $1/j$ as $j \rightarrow \infty$). The implication of these results is that the conditional variance of Y given X depends on the local-to-zero spectrum, with the shape of S for, say, $\omega < 12\pi$, essentially determining its value, even for large q . In terms of the original time series, frequencies of $|\omega| < 12\pi$ correspond to cycles of periodicity $T/6$. For instance, with 60 years worth of data (of any sampling frequency), the shape of the spectrum for frequencies below 10 year cycles essentially determines the uncertainty of the forecast of mean growth over the next 30 years.

Figure 2: Integral weights on local-to-zero spectrum for selected elements of Σ



Notes: The curves denote weights for variances (bold) and covariances (thin) of $(X_1, X_2, X_3, X_{10}, X_{11}, X_{12}, Y)$ with $w_{j,q+1}$ corresponding to $E(Y^2)$, $w_{q+1,q+1}$ corresponding to $E(Y^2)$, $w_{j,q+1}$ corresponding to $E(X_j Y)$, and $w_{j,k}$ on covariance terms $\Sigma_{jk} = E(X_j X_k)$ with $j, k \leq q$ and $j+k$ odd are equal to zero.

3.2 Prediction Sets in the $I(d)$ Model

A leading example of this analysis is given by the fractional $I(d)$ model, which has a (pseudo-) spectrum proportional to $|\lambda|^{-2d}$ for λ close to zero; this yields the local-to-zero spectrum $S(\omega) \propto |\omega|^{-2d}$, and the central limit result from the last subsection is applicable for $-1/2 < d < 3/2$. The $I(d)$ model captures a wide range of long-run dependence patterns including the usual $I(0)$ and $I(1)$ models, but also persistence patterns between and outside these two extremes. With negative values of d it also allows for long-run anti-persistence (which may arise from overdifferencing), and with $d > 1$ it allows for processes more persistent than an $I(1)$ process.

Running example (continued): Panels (i) of Figure 3 shows the appropriately centered and scaled Student-t predictive densities from (6) for the average growth rate of U.S. real per-capita GDP and the average value of CPI-inflation over the next 25 years for various values of d in the $I(d)$ model. For real GDP growth rates, predictive densities are shown for $d = -0.4, 0.0, 0.2$ and 0.5 , and for inflation the predictive densities are shown for $d = 0.0, 0.4, 0.7$, and 1.0 . For both series, as d increases, the variance of the predictive density increases because more persistence leads to larger variability in future average growth. The mode of the $I(0)$ predictive density is given by the in-sample mean (see the discussion following equation (6)), and the mode shifts to the left for $d > 0$ reflecting the persistent effect of the slow growth and low inflation experienced at the end of sample. In contrast, the mode of the $d = -0.4$ predictive density (shown for real GDP growth rates) is larger than the in-sample mean because faster than average growth is required to return the log-level of GDP to its pre-Great Recession trend growth path.

Evidently, both the length and location of 25-year ahead prediction sets depend critically on the d . This raises the question: What is the value of d for these series?

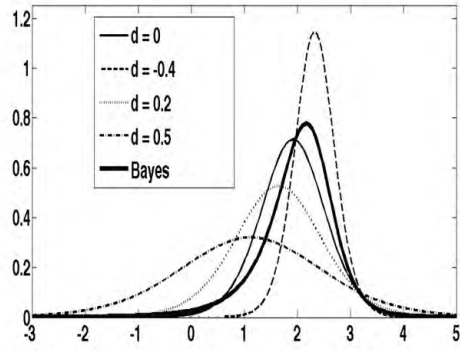
Panels (ii) summarize what the sample data say about the value of d . It plots the “low-frequency” log-likelihood values for d based on $X_{T,1:12}^s$ and its large-sample density from (5), and with the log-likelihood of the $I(0)$ model normalized to zero. The numbers for real per-capita GDP suggest only limited persistence for this series (values of $d > 0.6$ yield a log-likelihood 3 points lower than the $I(0)$ model), but values of d ranging from -0.4 (suggesting some reversion to a linear trend in the log-level of GDP, so that the growth rate is overdifferenced) to 0.2 (slight persistence in the GDP growth rates) all fit the data reasonably well. In contrast the inflation data suggest much more persistence: the log-likelihood has a maximum around $d = 0.6$ with corresponding log-likelihood values that are between 1.9 and 2.9 points larger than the $I(0)$ model.

Taken together, the results in panels (i) and (ii) indicate that much of the 25-year-ahead

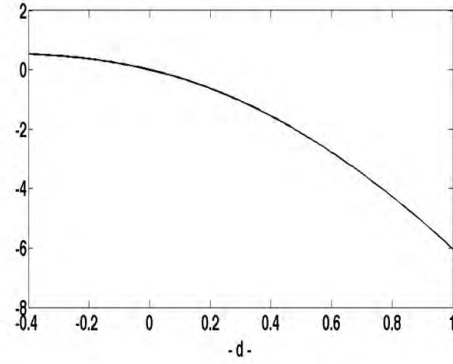
Figure 3: Predictive densities and low-frequency log likelihood values for the $I(d)$ model

(a) Growth rate of real per-capita GDP

(i) 25-year ahead $I(d)$ -predictive densities

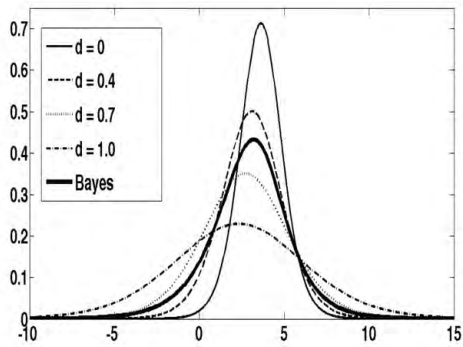


(ii) Low-frequency log-likelihood values for d

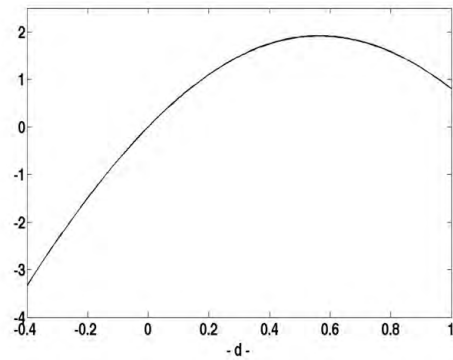


(b) Inflation (CPI)

(i) 25-year ahead $I(d)$ -predictive densities



(ii) Low-frequency log-likelihood values for d



Notes: Panels (i) show the known- d prediction sets and Bayes prediction sets using the prior $d \sim U[-0.4, 1.0]$. The low-frequency $l(d)$ likelihood is computed using $X_{T,1q}^d$ and its asymptotic distribution given in the text; values are relative to the $I(0)$ model.

forecast uncertainty is associated with uncertainty about the degree of persistence in the stochastic process, which in the $I(d)$ model is governed by the value of the parameter d . $\blacktriangle$

Bayes Prediction Sets: A natural way to incorporate this parameter uncertainty is to use a Bayes approach, where the limited sample information is combined with a prior on d . This is straightforward: With Γ the prior on d , the Bayes predictive density for Y^s conditional on $X^s = x^s$ is given by

$$f_{Y^s|X^s}^\Gamma(y^s|x^s) = \frac{\int f_{(X^s,Y^s)|d}(x^s,y^s)d\Gamma(d)}{\int f_{X^s|d}(x^s)d\Gamma(d)}$$

with $f_{(X^s,Y^s)|d}$ and $f_{X^s|d}$ the densities of (X^s, Y^s) and X^s in (5) with the value of Σ implied by a local-to-zero spectrum $S(\omega)$ proportional to $|\omega|^{-2d}$.

Bayes prediction sets can be readily computed from the predictive density. For example the “highest predictive density” (HPD) set for Y^s is $A^{HPD}(x^s) = \{y^s : f_{Y^s|X^s}^\Gamma(y^s|x^s) > cv(x^s)\}$, where $cv(x^s)$ solves $\int_{A^{HPD}(x^s)} f_{Y^s|X^s}^\Gamma(y^s|x^s)dy^s = 1 - \alpha$. This HPD Bayes set is the smallest length set that satisfies the coverage constraint relative to $f_{Y^s|X^s}^\Gamma$. Alternative Bayes prediction sets, such as equal-tailed sets, can be used instead. Thus, let $A^{\text{Bayes}}(x^s)$ denote a generic Bayes prediction set for Y^s as a function of x^s . Because $Y^s = Y/\sqrt{x'x}$ and $x^s = x/\sqrt{x'x}$, equivariance implies the extension to generic x via $A^{\text{Bayes}}(x) = \{y : y/\sqrt{x'x} \in A^{\text{Bayes}}(x/\sqrt{x'x})\}$.

Running example (continued): Panels (i) of Figure 3 shows the resulting Bayes predictive densities for $\bar{x}_{T:T+h}$ with a uniform prior on $d \in [-0.4, 1.0]$. This mixture of Student-t densities is no longer necessarily symmetric, as the underlying Student-t densities don’t have the same mode. So for instance, for the GDP series, one obtains a left-skewed Bayes predictive distribution since larger values of d both increase uncertainty and shift the most likely future values to the left. $\blacktriangle$

3.3 Frequentist Robustification of Bayes Prediction Sets

As discussed in Section 3.1, the distributions of (X, Y) and (X^s, Y^s) depend on the covariance matrix Σ , which in turn depends on the low-frequency spectrum S of x_t . In the next section, we discuss a parameterization of the spectrum that is more general than the $I(d)$ model, so in general, $\Sigma = \Sigma(\theta)$ where θ is a parameter vector. In this section we discuss the general problem of constructing frequentist prediction sets that incorporate uncertainty about the value of θ . We provide additional details in Appendix 8.2-8.4.

The (frequentist) coverage probability of a set A , $P_\theta(Y \in A(X))$, generally depends on the value θ . A Bayes prediction set has coverage probability of $1 - \alpha$, *on average* relative to the prior Γ , that is $\int P_\theta(Y \in A^{\text{Bayes}}(X))d\Gamma(\theta) = 1 - \alpha$, but in general, $P_\theta(Y \in A^{\text{Bayes}}(X)) < 1 - \alpha$

for some values of θ . In this subsection, we "robustify" Bayes sets by enlarging them so they have frequentist coverage: $\inf_{\theta \in \Theta} P_{\theta}(Y \in A(X)) \geq 1 - \alpha$. There is no unique way to achieve this. We focus on sets with smallest weighted expected length.

To be specific, let $A(X)$ denote an arbitrary prediction set, and $V_{\theta}(A) = E_{\theta}[\text{vol}(A(X))]$ denote its expected length (which depends on θ). The goal is to choose A to minimize $V_{\theta}(A)$ over the parameter space Θ for θ . In many problems, including the one considered in this paper, there is no set A that simultaneously minimizes $V_{\theta}(A)$ for all $\theta \in \Theta$ while maintaining coverage, so there is an inherent trade-off of expected length over different values of θ . Let W denote a weighting function that makes this trade-off explicit. Consider the following problem:

$$\min_A \int V_{\theta}(A) dW(\theta) \quad (8)$$

subject to

$$\text{Equivariance:} \quad y \in A(x) \text{ implies } by \in A(bx) \text{ for all } x, y \text{ and } |b| \neq 0 \quad (9)$$

$$\text{Frequentist Coverage:} \quad \inf_{\theta \in \Theta} P_{\theta}(Y \in A(X)) \geq 1 - \alpha, \text{ and} \quad (10)$$

$$\text{Bayes Superset:} \quad A^{\text{Bayes}}(x) \subset A(x) \text{ for all } x. \quad (11)$$

Because the objective function depends on the weighting function W , so will the solution, and we discuss specific choices for W in the following section. The constraint (9) imposes scale invariance – recall that location invariance in the original problem is imposed by the choice of Y and X . The coverage constraint that defines a $(1 - \alpha)$ -frequentist prediction set is given by (10).

The constraint (11) can be motivated in a variety of ways. One motivation is *ad hoc* and simply says that the goal is to robustify a Bayes set by enlarging it so it has frequentist coverage properties. Another focuses on properties of frequentist sets that do not impose (11). Notably, conditional on particular realizations of X these sets can have unreasonably small length; indeed they can be empty. In particular, even with θ known (i.e., $\Theta = \{\theta\}$), solving (8) subject to (9) and (10) does *not* in general yield the known- θ prediction set (6), but rather a prediction set whose coverage of Y is equal to $1 - \alpha$ only on average over repeated draws of X , but not conditional on the observed X . Müller and Norets (2012) show that imposing (11) eliminates these arguably unattractive properties. We find the Müller and Norets arguments compelling and therefore enforce the constraint (11) for the frequentist sets used in the empirical analysis of Section 6. However, for comparison we also study solutions that do not impose (11) in Section 4.

The solution to the program (8)-(11) can be found in three steps: the first step transforms the problem to impose equivariance (9); the second uses a "least favorable distribution" for θ

to simplify the coverage constraint (10); and the third enforces (11). We discuss these steps in turn.

Equivariance: If $A(X)$ is scale equivariant, then $Y \in A(X)$ if and only if $Y \in \sqrt{X'X}A(X^s)$. Thus, $\text{vol}(A(X)) = \sqrt{X'X} \text{vol}(A(X^s))$ and $V_\theta(A) = E_\theta [g_\theta(X^s) \text{vol}(A(X^s))]$, where $g_\theta(X^s) = E_\theta[\sqrt{X'X}|X^s]$. Imposing this restriction, the objective function (8) becomes

$$\min_A \int E_\theta [g_\theta(X^s) \text{vol}(A(X^s))] dW(\theta), \quad (12)$$

and the coverage (10) and Bayes superset (11) constraints can be rewritten as

$$\inf_{\theta \in \Theta} P_\theta(Y^s \in A(X^s)) \geq 1 - \alpha \quad (13)$$

$$A^{\text{Bayes}}(x^s) \subset A(x^s) \text{ for all } x^s. \quad (14)$$

Note that (12)-(14) only involve the value of A evaluated at x^s , which lives on a smaller subspace $x^{s'}x^s = 1$ compared to $x \in \mathbb{R}^q$, but on that subspace, A is unrestricted. The solution to (12) subject to (13) and (14), $A^*(x^s)$, then implies the solution $A^*(x) = \{y : y/\sqrt{x'x} \in A^*(x/\sqrt{x'x})\}$ to the original problem (8) subject to (9)-(11).

Frequentist Coverage: For the coverage constraint (13), suppose for a moment that θ is a random variable with distribution Λ , and consider solving (12) subject to the resulting single coverage constraint

$$\int P_\theta(Y^s \in A(X^s)) d\Lambda(\theta) \geq 1 - \alpha. \quad (15)$$

A calculations yields the solution

$$A_\Lambda(x^s) = \left\{ y^s : \frac{\int f_{(Y^s, X^s)|\theta}(y^s, x^s) d\Lambda(\theta)}{\int g_\theta(x^s) f_{X^s|\theta}(x^s) dW(\theta)} > \text{cv} \right\} \quad (16)$$

where cv is chosen to satisfy (15) with equality. Of course, while A_Λ satisfies the *average* coverage constraint (15), it does not necessarily satisfy the *uniform* coverage constraint (13) required for a frequentist prediction set. However, because any set satisfying (13) also satisfies (15), the value of the objective (12) evaluated at A_Λ provides a lower bound for any set satisfying (13). Therefore, *if* a distribution $\Lambda^\dagger$ can be found so that $A_{\Lambda^\dagger}$ satisfies (13), it solves the minimization problem (12) subject to the uniform coverage constraint in (13). Such a $\Lambda^\dagger$ is called the “least favorable distribution” for the problem. Elliott, Müller, and Watson (2014) develop numerical methods for approximating least favorable distributions in related problems, and we use a variant of those methods here.

Bayes Superset: The final step – incorporating the constraint (14) – is straightforward: it simply amounts to replacing (16) with the set

$$A^{MN}(x^s) = \left\{ y^s : \frac{\int f_{(Y^s, X^s)|\theta}(y^s, x^s) d\Lambda^\dagger(\theta)}{\int g_\theta(x^s) f_{X^s|\theta}(x^s) dW(\theta)} > cv^{MN} \right\} \cup A^{\text{Bayes}} \quad (17)$$

where $(\Lambda^\dagger, cv^{MN})$ are now such that $\int P_\theta(Y^s \in A^{MN}(X^s)) d\Lambda^\dagger(\theta) = 1 - \alpha$ and $\inf_{\theta \in \Theta} P_\theta(Y^s \in A^{MN}(X^s)) \geq 1 - \alpha$ (cf. Theorem 4 in Müller and Norets (2012)).

4 Parameterizations for Long-Horizon Prediction Sets

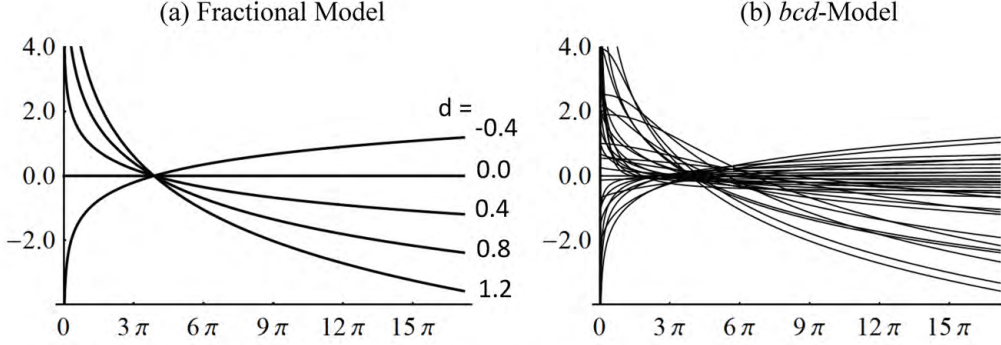
Implementation of the prediction sets discussed in the last section requires four ingredients: (i) a parameterization of S , the local-to-zero spectrum, which yields the covariance matrix $\Sigma(\theta)$ via (7); (ii) a Bayes prior $\Gamma(\theta)$, which yields the Bayes prediction set A^{Bayes} ; (iii) a frequentist weighting function $W(\theta)$, which quantifies the trade-off of expected length for various of θ in the objective function (8); and (iv) a choice for q , the number of cosine weighted averages used for the prediction problem. These are discussed in the following three subsections.

4.1 Parameterizing the Low-Frequency Spectrum

The $I(d)$ model introduced in Section 3.2 above is a flexible one-parameter model that captures a wide range of long-run persistence patterns. Because of its simplicity, flexibility, and use in other empirical analyses involving long-run behavior of economic time series, we use the $I(d)$ parameterization for our equal-tailed Bayes prediction sets A^{Bayes} .

However, a concern is that the family of $I(d)$ local-to-zero spectra may not be sufficiently flexible to capture all forms of long-run dependencies in economic time series. This suggests the need for a richer class of local-to-zero spectra, and we construct such a class by considering two other models that have proven useful for modelling low-frequency characteristics in other contexts. The first is the local-level model which expresses x_t as the sum of an $I(0)$ process and an $I(1)$ process, say $x_t = e_{1t} + (bT)^{-1} \sum_{s=1}^t e_{2s}$, where $\{e_{1t}\}$ and $\{e_{2t}\}$ are mutually uncorrelated $I(0)$ processes with the same long-run variance. The $I(1)$ component has relative magnitude $1/b$ and is usefully thought of as a stochastically varying ‘local mean’ of the growth rate x_t , as arising from some forms of stochastic breaks. In this model, $S(\omega) \propto b^2 + \omega^{-2}$. The second model is the local-to-unity AR(1) model, widely used to model highly persistent processes. In this model $x_t = (1 - c/T)x_{t-1} + e_t$, where e_t is an $I(0)$ process,

Figure 4: Logarithm of the local-to-zero spectrum for selected models



and a straightforward calculation shows that $S(\omega) \propto 1/(\omega^2 + c^2)$. (Note that $(b, c) \rightarrow (\infty, \infty)$ and $(b, c) \rightarrow (0, 0)$ recover the $I(0)$ and $I(1)$ model, respectively). The $I(d)$, local-level and local-to-unity models are nested in the parameterization

$$S(\omega) \propto \left(\frac{1}{\omega^2 + c^2} \right)^d + b^2 \quad (18)$$

where $b = c = 0$ for the $I(d)$ model, $d = 1$, and $c = 0$ for the local-level model, and $d = 1$, and $b = 0$ for the local-to-unity model.⁹

Figure 4 plots the logarithm of the local-to-zero spectrum of the $I(d)$ model in panel (a), and of this bcd -model in panel (b). The bcd -parameterization allows us to capture a wide range of monotone shapes for the low frequency (pseudo-) spectrum of x_t , including, but not limited to, the three benchmark models discussed above. In the analysis below we let $\theta = (b, c, d)$, so that $\Sigma(\theta)$ is given by (7) with the local-to-zero spectrum S as in (18).

4.2 Bayes and Frequentist Weighting Functions

In the empirical analysis in Section 6, we assume that S is characterized by the bcd -model, with $-0.4 \leq d \leq 1.0$ and $b, c \geq 0$.¹⁰ As mentioned above, we construct Bayes sets using a

⁹This is recognized as the local-to-zero spectrum of the process $x_t = e_{1t} + (bT^d)^{-1}z_t$, where $(1 - \rho_T L)^d z_t = e_{2t}$ with $\rho_T = 1 - c/T$ and $\{e_{1t}\}$ and $\{e_{2t}\}$ are mutually uncorrelated $I(0)$ processes with the same long-run variance. It is also recognized as the spectrum of the Whittle-Matérn process from spatial statistics (e.g., Lindgren (2013)). Autocovariances for this process are derived in the appendix.

¹⁰In an earlier version of this paper we allowed for values of d as large as 1.4. However upon reflection, values of d larger than 1.0 seem unnecessary for the economic variables we study (growth rates of real variables, inflation rates, and asset returns). Larger values of d may be appropriate in other applications,

prior that puts all weight on the $I(d)$ model (so that $b = c = 0$); we use a prior with uniform weight on values of $d \in [-0.4, 1.0]$. The A^{MN} sets robustify these Bayes sets so they have frequentist coverage for all values of $b, c \geq 0$ and $d \in [-0.4, 1.0]$. The analysis is usefully thought of in terms of the various spectral shapes plotted in Figure 4, and the Bayes prior is seen as putting equal weight on the various shapes in panel (a). Because S may take on shapes other than those represented by the $I(d)$ models in panel (a), the A^{MN} sets robustify the Bayes analysis to ensure frequentist coverage over all shapes shown in panel (b).

Construction of the A^{MN} sets requires specification of the weighting function W in (8). As noted in Section 3.3, the function W determines the trade-off between expected length for various of θ , which is necessary because there is no single prediction set that minimizes expected length for all θ . Our choice of W is guided by the observation that, even with θ known, the minimized values of $V_\theta(A)$ vary greatly over the values of θ . For example, in the $I(d)$ model with known d , prediction sets are much wider when $d = 1$ (so that $x_t \sim I(1)$) than when $d = 0$ ($x_t \sim I(0)$). To account for these differences we scale $V_\theta(A)$ so that it is expressed in units of the expected length of the predictions set for known θ . Denote this scaled version of $V_\theta(A)$ by $R_\theta(A) = V_\theta(A)/V_\theta^{\text{known}}$, where V_θ^{known} is the expected length of the prediction set for known value of θ implied by (6). This relative measure of length is readily interpretable: $R_\theta(A) > 1$ represents the "regret" about not knowing θ when using the set A . In terms of $R_\theta(A)$ we use a weighting function that coincides with the Bayes prior: uniform on $d \in [-0.4, 1.0]$ and with $b = c = 0$ (so in terms of $V_\theta(A)$, the weighting function W is proportional to $1/V_\theta^{\text{known}}$). Thus, we robustify the Bayes sets in a way that minimizes regret using the Bayes prior while achieving frequentist coverage. We investigate how this weighting function performs relative to other possible weighting functions below, but first we consider the coverage rates of the various sets.

Table 1 shows coverage rates for 67% and 90% prediction sets for $h = rT$, with $r = 0.5$ using $q = 12$ cosine transforms. (This is the value of q we will use in the empirical analysis, and is discussed more fully in the next subsection). Table 1 answers two questions. First, what is the frequentist coverage of the Bayes prediction sets across the range of processes represented by the spectra in panels (a) and (b) of Figure 4? And second, does the $I(d)$ model provide sufficient flexibility so that the additional parameters b and c are unnecessary in practice? The table therefore displays coverage rates for three prediction sets: the Bayes set, A^{Bayes} ; the set robustified to have correct frequentist coverage over d but with $b = c = 0$, denoted A_d^{MN} ; and the set robustified to have correct frequentist coverage over (b, c, d) , $A_{(b,c,d)}^{MN}$. Coverage rates are shown for three configurations of (b, c, d) . In the first, values of (b, c, d) are drawn from the prior, so A^{Bayes} has correct coverage; in the second, the coverage

and we note that the results in Section 3.1 hold for the bcd -model with $-0.5 < d < 1.5$, and $b, c \geq 0$.

Table 1: Coverage for nominal 67% and 90% prediction sets, $r = 0.5, q = 12$

	67% Prediction Sets				90% Prediction Sets		
	A^{Bayes}	A_d^{MN}	$A_{(b,c,d)}^{MN}$		A^{Bayes}	A_d^{MN}	$A_{(b,c,d)}^{MN}$
$b = 0, c = 0, d \sim U[-0.4, 1.0]$	0.67	0.71	0.73		0.90	0.93	0.94
<i>Coverage minimized over:</i>							
$b = 0, c = 0, -0.4 \leq d \leq 1.0$	0.55	0.67	0.68		0.81	0.90	0.90
$b \geq 0, c \geq 0, -0.4 \leq d \leq 1.0$	0.54	0.57	0.67		0.81	0.84	0.90

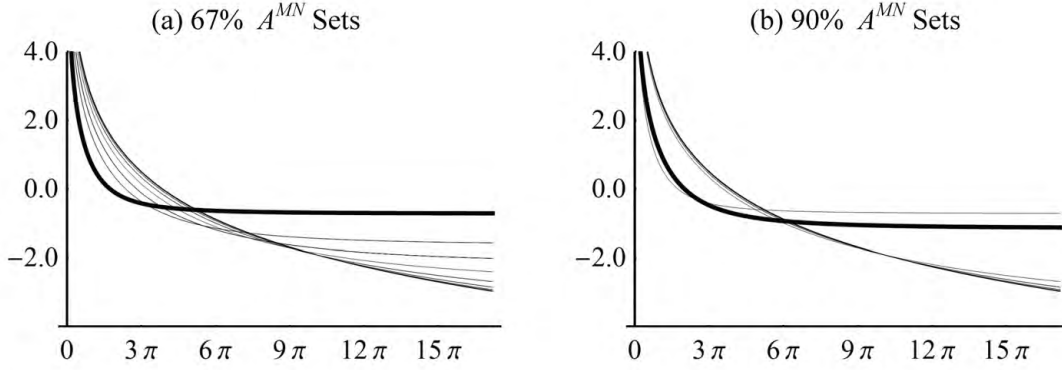
Notes: A^{Bayes} is the equal-tailed Bayes prediction set using the prior $b = 0, c = 0, d \sim U[-0.4, 1.0]$; A_d^{MN} robustifies the Bayes set so it has frequentist coverage for $b = 0, c = 0, d \in [-0.4, 1.0]$; $A_{(b,c,d)}^{MN}$ robustifies the Bayes set so it has frequentist coverage for $b \geq 0, c \geq 0, d \in [-0.4, 1.0]$.

probability is minimized over $-0.4 \leq d \leq 1.0$ with $b = c = 0$, so A_d^{MN} has correct coverage; and in the third, the coverage probability is also minimized over $b, c \geq 0$, so $A_{(b,c,d)}^{MN}$ has the correct coverage. The table indicates that A^{Bayes} exhibits substantial undercoverage for some values of d and (b, c, d) . It also indicates substantial undercoverage of A_d^{MN} for some values of (b, c, d) . Evidently, controlling coverage over d does not provide adequate coverage for long-run persistence patterns associated with non-zero values of b and c . Thus, because some economic variables are arguably well-described by stochastic processes with non-zero value of b and c , it seems prudent to construct the $A_{(b,c,d)}^{MN}$ sets.

Figure 5 displays the approximate least favorable distributions (ALFDs) that underlie the $A_{(b,c,d)}^{MN}$ sets. Were the A_d^{MN} sets adequate, these ALFDs would put no mass on non-zero values of b and c . We display the ALFDs in terms of the local-to-zero spectral shapes familiar from Figure 4, with thicker curves representing more mass in the ALFD. The ALFD is non-degenerate and has most of its mass on spectra that are relatively flat for larger ω , but with a pronounced pole at zero (these spectra arise, for instance, in the local-level with moderate b). Intuitively, in the local-level model, the strong mean reversion of the $I(0)$ component masks the pronounced long-run uncertainty, making it relatively hardest to control coverage.

Table 1 investigated the coverage rates of the Bayes and A^{MN} sets for a benchmark weighting function W and prior Γ . Table 2 examines the expected length of the $A^{MN} = A_{(b,c,d)}^{MN}$ sets and answers questions regarding the choice of W and Γ . Because A^{MN} minimizes a weighted average of the values of the expected (normalized) length $R_\theta(A)$, it is impossible to choose an alternative weighting function with uniformly lower values of $R_\theta(A)$. However, it is possible to reduce $R_\theta(A^{MN})$ for any particular value of θ . Thus, Table 2 compares the value $R_\theta(A^{MN})$ for the benchmark choice of W to the value obtained by putting all W mass

Figure 5: ALFDs in terms of implied local-to-zero spectra



Notes: The panels represent the logarithm of the local-to-zero spectra that receive more than 1% mass in the approximate least favorable distribution underlying the A^{MN} predictive sets for $r=0.5$ and $q=12$. Line widths are proportional to mass in the ALFD. Same scale normalization as in Figure 4.

Table 2: Relative volume, $R_\theta(A)$, for 90% A^{MN} prediction sets, $r = 0.5$, $q = 12$

	Parameter values (θ)											
b	0	0	0	0	0	0	2	0.2	0.04	0	0	0
c	0	0	0	0	0	0	0	0	0	0.82	3.32	27.11
d	-0.40	-0.20	0.00	0.40	0.80	1.00	1	1	1	1	1	1
(a) Sets constrained to include Bayes set												
W^{Bench}	1.85	1.88	1.80	1.54	1.27	1.15	1.78	1.38	1.22	1.89	1.19	1.40
Env	1.80	1.85	1.76	1.47	1.22	1.02	1.73	1.31	1.12	1.76	1.09	1.32
(b) Sets unconstrained												
W^{Bench}	1.69	1.84	1.80	1.54	1.27	1.15	1.77	1.38	1.22	1.88	1.19	1.40
Env	1.57	1.79	1.74	1.47	1.22	1.02	1.71	1.30	1.12	1.67	1.09	1.32

Notes: The table shows the values of $R_\theta(A)$ for the value of $\theta = (b, c, d)$ shown in the first three rows. W^{Bench} denotes the benchmark weighting function described in the text, and in the rows labeled " W^{Bench} ," A minimizes a weighted average of $R_\theta(A)$; in the rows labeled "Env," A is a θ -specific set that minimizes expected length at that value of θ . The non-zero values of b and c are chosen to lie between the $I(0)$ ($b \rightarrow \infty$, $c \rightarrow \infty$) and $I(1)$ ($b=0$, $c=0$) extremes.

on θ , for selected values of θ . Said differently, Table 2 compares $R_\theta(A^{MN})$ to the envelope obtained by minimizing $R_\theta(A)$ over all equivariant sets A that have uniform coverage in the bcd -model. Values of $R_\theta(A)$ are shown as a function of d in the $I(d)$ model, and for selected values in the b and c models, respectively.

The table shows results for 90% prediction sets; conclusions for 67% prediction sets are similar. Panel (a) considers sets that are also restricted to contain the Bayes prediction set, as in (11). For the benchmark choice of W , $R_\theta(A) = R_\theta(A^{MN})$ ranges from 1.15 (so that $V_\theta(A^{MN})$ is 15% larger than the attainable expected length with θ known) when $b = c = 0$ and $d = 1.0$, to 1.89 when $b = 0$, $c = 27$, and $d = 1.0$. These numbers demonstrate that parameter uncertainty is a significant component of A^{MN} . The envelope numbers are smaller, of course, but are typically within 10% of the A^{MN} values even for θ with non-zero b or c (where the benchmark weighting function puts zero mass).

Panel (b) of Table 2 takes up the question of the cost of imposing the Bayes-superset constraint (11). The two rows in panel (b) are computed exactly as those in panel (a), except that the Bayes-superset constraint is no longer enforced. Evidently, for most values of θ , the Bayes-superset constraint has a negligible impact of the expected length of the prediction set. What is more, the envelope numbers in panel (b) are typically not much smaller than those for A^{MN} in panel (a), which means that the substantial regret of not knowing θ is mostly driven by the coverage constraint, and is not an artifact of our benchmark choices for Γ or W . (And, as a corollary, it also implies that alternative choices for Γ for A^{MN} can not lead to substantially shorter sets on average).

We draw four conclusions from Tables 1 and 2. First, Bayes sets constructed using a uniform prior on d exhibit substantial undercoverage for some values of d . Second, robustifying these sets to achieve frequentist coverage over d is inadequate for some processes with non-zero values of b and c . Third, for many values of (b, c, d) our benchmark choices of Γ and W produce sets with expected length close to the smallest achievable length under the coverage constraint. And finally, for most values of (b, c, d) there is little cost in terms of expected length for constructing frequentist sets that are supersets of Bayes sets (and therefore share some their desirable properties).

4.3 Choice of q

As discussed in Section 2, the choice of q may usefully be thought of as a trade-off between efficiency and robustness. In principle, the central limit theorem for $(X'_{T1:q}, Y_T)'$ discussed in Section 3.1 holds for any fixed q , at least asymptotically. And the larger q , the smaller the (average) uncertainty about Y_T . This suggests that one should pick q large to increase

Table 3: Average length of 90% A^{MN} prediction sets, $r = 0.5$

	$q = 12$		$q = 6$	$q = 24$	$q = 48$
θ unknown	1.57		1.69	1.45	1.38
θ known	1.00		1.08	0.97	0.95

Notes: Entries are weighted average regret (relative to knowing θ in the $q=12$ case) of A^{MN} prediction sets (θ unknown) and minimum length sets A_θ^{known} for θ known.

efficiency of the procedure.

At the same time, one might worry that approximations provided by the central limit theorem for $(X'_{T,1:q}, Y_T)'$ become poor for large q . The concern is not only that the high-dimensional multivariate Gaussianity might fail to be an accurate approximation; more importantly, any parametric assumption about the shape of the local-to-zero spectrum becomes stronger for larger q . In particular, for a given sample size T , the assumption that the spectrum of x_t over the frequencies $[-q\pi/T, q\pi/T]$ is well approximated by the spectrum of the *bcd*-model becomes less plausible the larger q . Roughly speaking, we fit a parametric model to the q observations $X_{T,1:q}$, so a concern about nontrivial approximation errors arises for large q , irrespective of the sample size T .

We are thus faced with a classic efficiency and robustness trade-off. Recall from the discussion of Figure 2 in Section 3.1, however, that the object of interest – the variability of long-run forecasts, as embodied by the conditional variance of Y given X – is a low frequency quantity that is essentially governed by properties of x_t over frequencies $[-12\pi/T, 12\pi/T]$. Since the predictors $X_T(j)$ provide information for frequency $j\pi/T$, this suggests that the marginal benefit of increasing q beyond $q = 12$ is modest, at least with the spectrum known.

With the spectrum unknown, X with larger q provides additional information about its scale and its shape. The scale effect is most easily understood in the $I(0)$ model. As discussed above, the $I(0)$ prediction set is $\bar{x}_{1:T} \pm t_{(1-a/2)}^q (1 + r^{-1})^{1/2} T^{-1/2} s_{LR}$, where $s_{LR}^2 = (T/q) X'_{T,1:q} X_{T,1:q}$. The average asymptotic length of this forecast is thus $2T^{-\alpha} t_{(1-a/2)}^q (1 + r^{-1})^{1/2} E\sqrt{X'X/q}$ with $X \sim \mathcal{N}(0, \sigma^2 I_q)$, which decreases in q , since $t_{(1-a/2)}^q E\sqrt{X'X/q}$ is a decreasing function of q .¹¹ But the benefit of increasing q is modest: for a 90% interval, the average length for $q \in \{24, 48, \infty\}$ is only $\{3.0\%, 4.4\%, 5.8\%\}$ shorter than for $q = 12$, for instance.

When the shape of the spectrum is unknown but parametrized, as in the *bcd*-model,

¹¹This is analogous to the wider confidence intervals that arise from the use of inconsistent HAC estimators; see Kiefer, Vogelsang, and Bunzel (2000) and Kiefer and Vogelsang (2002, 2005), and Müller (2014) for a review.

Table 4: 25-year-ahead prediction sets

Series	67%				90%		
	A^{Bayes}	A^{MN}	$A^{I(0)}$		A^{Bayes}	A^{MN}	$A^{I(0)}$
GDP/Pop	(1.5, 2.7)	(1.5, 2.7)	(1.3, 2.5)		(0.8, 3.0)	(0.5, 3.0)	(0.8, 3.0)
Inflation	(0.7, 5.1)	(0.7, 5.1)	(2.4, 4.9)		(-1.4, 6.8)	(-2.2, 7.1)	(1.5, 5.8)

Notes: The table shows the prediction sets for the average value of the growth rate of real per-capita GDP and inflation from 2012-2037 using the benchmark values of the Γ , W , and q .

increasing q beyond 12 provides additional information about the shape of the spectrum over the crucial frequencies $[-12\pi/T, 12\pi/T]$. Table 3 quantifies the combined scale and shape effects by reporting the value of the objective $\int V_\theta(A)dW(\theta)$ in the program (8) for $q \in \{6, 12, 24, 48\}$. To be able to make meaningful comparisons across different values of q , we choose W as the benchmark weighting function from the $q = 12$ case for all values of q , normalized to integrate to one (so that the values in Table 3 are the weighted average regret relative to knowing θ in the $q = 12$ case). In this θ unknown case, there is an 8% decrease in average length as q increases from $q = 12$ to $q = 24$ and a further reduction of 5% for $q = 48$. As a point of reference, the table also reports the same average length for θ known.¹² The reductions in lengths for $q > 12$ are only marginally larger than the pure scale effect in the $I(0)$ model, corroborating the intuition above about the relative unimportance of the spectral shape outside the $[-12\pi/T, 12\pi/T]$ interval for long-run forecasting.

Much of our empirical analysis uses 65 years of post-WWII data, so that a choice $q = 12$ relies on periodicities below 10.8 years, while $q = 24$ and $q = 48$ use periodicities below 5.4 and 2.7 years. The marginal benefit of increasing q beyond $q = 12$ is modest, and q must be chosen very much larger to substantially reduce forecast uncertainty overall. In our view, a concern about severe spectral misspecification outweighs these potential gains, so that as a default, we suggest constructing the prediction sets with $q = 12$.

Running example (continued): Table 4 shows the 67% and 90% A^{Bayes} and A^{MN} 25-year-ahead predictions sets for real GDP growth and inflation using the benchmark values of the Bayes prior (Γ), weighting function (W), and $q = 12$. The 67% A^{Bayes} and A^{MN} sets coincide, while the 90% A^{MN} sets are somewhat wider than the A^{Bayes} sets. For comparison, the table also shows the prediction sets computed from the $I(0)$ model. These are similar to the A^{Bayes} and A^{MN} sets for GDP (although the 67% $I(0)$ set is shifted to the left for reasons discussed above), but are much different for inflation (where the $I(0)$ are shifted the

¹²These are for comparison only; the parameters of the *bcd*-model cannot be consistently estimated even under $q \rightarrow \infty$ asymptotics.

right and are much narrower), and where both results are as expected given the predictive densities and log-likelihood values displayed in Figure 3. Section 6 discusses these empirical results in more detail. ▲

5 Finite Sample Experiments

In the last two sections we developed a large-sample framework for constructing Bayes and frequentist long-run prediction sets that is tailored to models of long-run persistence typically used for economic time series. This large sample analysis is sufficiently general to allow for in-sample and out-of-sample stochastic breaks in the series, as long as these breaks occur with sufficient frequency that sample averages satisfy the central limit theorem discussed in Section 3. And the large-sample analysis also accommodates short memory stochastic shifts in volatility. But does this large-sample analysis provide reliable prediction sets for the sample sizes and stochastic processes typically encountered in applied economics? This section addresses this question using two sets of finite sample experiments. The first set of experiments are Monte Carlo simulations in which we generate data with level and volatility breaks designed to mimic the kinds of breaks seen in some macroeconomic time series. The second set of experiments use a rolling sample of daily returns and squared returns of the SP500 index to construct pseudo-out-of-sample prediction sets and uses actual values of returns to evaluate these sets. We discuss these experiments in the following two subsections.

5.1 Monte Carlo Simulations with Breaks in Level and Volatility

A concern with the methods developed in the previous sections is that they might not sufficiently account for the kinds of “breaks” or “shifts” in stochastic processes that sometimes occur with economic data. It is obviously true that arbitrary breaks can undermine any attempt at prediction, so it is correct – but uninteresting – to point out that the methods proposed here are not immune to arbitrarily defined breaks. However, it is interesting to ask how well the methods fare in the face of breaks that plausibly have occurred in the kinds of series to which the methods are to be applied, and that question is addressed in this subsection. Statistical characterizations of uncertainty require a probability framework, and therefore we consider breaks that occur probabilistically. And, because of the macroeconomic applications we carry out in Section 6, the models for these breaks are motivated by the behavior of important macroeconomic time economic series in the post-WWII United States.

We consider four models. The first two involve breaks in the level of x_t

$$x_t = \mu_t + u_t \quad (19)$$

where μ_t denotes the "level" of x_t and u_t is a zero-mean stochastic process that is independent of μ_t . We suppose that μ_t shifts discretely by an amount $\pm\delta$ at irregular time periods determined by the indicator s_t . That is

$$\mu_t = \mu_{t-1} + s_t\delta_t \quad (20)$$

where s_t is an i.i.d. Bernoulli process with $P(s_t = 1) = p$, and $\delta_t = \pm\delta$ with equal probability independent of s_t . Note that μ_t follows a martingale – an $I(1)$ process – so that in large samples its sample averages are characterized by the Gaussian limits in Section 3 (as an $I(1)$ model for fixed $p, \delta > 0$ and a special case of the “ b ” model in subsection 4.1 for fixed $p > 0$ and $\delta = O(T^{-1})$). That said, when p is small, shifts in μ_t occur infrequently and the finite sample behavior of sample averages may be quite different from their large-sample Gaussian limit.

The second two models involve breaks in volatility. In these models x_t has components that can be represented as $\sigma_t e_t$, where e_t is an $I(0)$ process and σ_t is a volatility process that evolves as $\ln(\sigma_t) = \mu_t$, where μ_t follows (20). While the central limit used in Section 3 allows for certain forms of heteroskedasticity, it does not allow volatility to evolve as an $I(1)$ process. Thus, the volatility models in this section involve stochastic processes that are strictly more general than the processes analyzed above, even in large samples.

We choose model parameters to match specific characteristics of post-WWII U.S. quarterly macroeconomic data. Thus, we set $T = 260$ (which corresponds to 65 years of quarterly observations), and as above we consider forecast horizons of $h = 0.5T = 130$ with $q = 12$ and the prior (Γ) and weighting function (W) described in Section 4. We choose two values for the break frequency: $p = 1/40$ (so a break occurs, on average, once every 40 quarters) and $p = 1/260$ (so a break occurs, on average, once during the sample period period). The other parameter values depend on the experiment and are motivated by the behavior of particular U.S. macroeconomic time series.

Model 1 is motivated by the growth rate of average labor productivity, which visually appears to be an $I(0)$ process but around a time varying level. (See appendix Figure A.4.) Labor productivity growth averaged 2.2% per year in the post-WWII period, but experienced decade-long swings that were roughly one percentage point higher (early 1960s and late 1990s) or lower (1970s and early 1980s) than the average. The first model therefore takes the form (19) with $u_t \sim iid\mathcal{N}(0, \sigma_u^2)$, where σ_u is chosen to match the long-run standard deviation of average labor productivity, and the magnitude of the breaks in μ_t was chosen

to yield a sensible value for the interquartile range (IQR) of $\mu_T - \mu_0$. Specifically, for each value of p we chose two values for δ , where the first, δ_{small} , yielded an IQR of 0.5% and the second, δ_{large} , yielded an IQR of 1.5%.

Model 2 is similar to Model 1, but is motivated by the behavior of nominal interest rates, which follow a pattern consistent with (19) but with u_t a highly serially correlated process. Thus for this experiment, u_t was generated by an AR(1) process with coefficient 0.98, Gaussian innovations with variance chosen to match 10-year U.S. Treasury Bonds, and δ_{small} and δ_{large} chosen so that the IQR for $\mu_T - \mu_0$ was 2.0% and 4.0%, respectively.

Model 3 is designed to capture features in the data like the "Great Moderation": a low-frequency reduction in the volatility in real U.S. macroeconomic variables. For example, the standard deviation of growth rates of measures of real aggregate activity (GDP, employment, etc.) fell rather abruptly by roughly 30% in the early 1980s (e.g., Stock and Watson (2002)). Thus, in this model the data were generated as $x_t = \sigma_t e_t$, with $e_t \sim iid\mathcal{N}(0, 1)$ and $\ln(\sigma_t) = \mu_t$ generated as described above with δ_{small} and δ_{large} chosen so that the IQR for $\ln(\sigma_T/\sigma_0)$ was 0.25% and 0.75%, respectively.

Model 4 is designed to capture the changes in variability and persistence evident in the U.S. inflation process. Cogley and Sargent (2014), Stock and Watson (2007), and others argue that these features can be captured by a local-level-model with stochastic volatility. Thus, in this model we generate data as $x_t = e_{1t} + \sum_{s=1}^t \sigma_s e_{2s}$ where e_{1t} and e_{2t} are mutually independent i.i.d. standard normal random variables, $\ln(\sigma_t) = \mu_t$ follows the process described above, and the parameters are chosen to mimic estimates of the time-varying $IMA(1,1)$ representation of the model found in U.S. data (e.g., Watson (2014)). Specifically, σ_0 is chosen so that MA coefficient is 0.5 in the initial period, and δ_{small} and δ_{large} were chosen so that the IQR of the full-sample change in the MA coefficient was 0.5 and 0.8.

Results for the various experiments are shown in Table 5, where panel (a) shows results for the A^{Bayes} sets and panel (b) shows results for the A^{MN} sets. The first row of each panel shows results for the model with $p = 0$ (so that breaks are absent); the other rows show results for $p = 1/40$, $p = 1/260$, and for δ_{small} and δ_{large} . When $p = 0$, models 1 and 3 are i.i.d. processes for which both A^{Bayes} and A^{MN} have coverage rates that exceed their nominal level. This "overcoverage" occurs because A^{Bayes} provides correct *average* coverage for $I(d)$ processes that includes both small and large values of d , and coverage for small d is less than the average coverage. Similar reasoning explains the overcoverage for A^{MN} , which is designed to achieve uniform coverage over (b, c, d) . And with $p = 0$, model 2 is well approximated by the local-to-unity model with $c = 260(1 - 0.98) = 5.2$ and model 4 is well approximated by an $I(1)$ process; A^{MN} satisfies the coverage constraint in both models, while A^{Bayes} severely undercovers in model 4, achieving the same undercoverage shown previously in Table 1 for

Table 5: Coverage probability for simulated data, $T = 260$, $r = 0.5$, and $q = 12$

		Model 1: Break in level, $I(0)$ model			Model 2: Break in level, AR(1) model			Model 3: Break in volatility, $I(0)$ model			Model 4: Break in volatility, $I(0)+I(1)$ model	
Nom. Level		67%	90%		67%	90%		67%	90%		67%	90%
		(a) A^{Bayes}										
$p = 0$		0.72	0.94		0.65	0.89		0.72	0.94		0.55	0.81
$p = 1/40$	δ_{small}	0.68	0.92		0.56	0.81		0.71	0.94		0.57	0.77
	δ_{large}	0.58	0.85		0.56	0.80		0.71	0.92		0.58	0.73
$p = 1/260$	δ_{small}	0.69	0.93		0.61	0.79		0.71	0.94		0.57	0.79
	δ_{large}	0.62	0.88		0.62	0.77		0.71	0.93		0.57	0.75
		(b) A^{MN}										
$p = 0$		0.74	0.95		0.70	0.94		0.74	0.95		0.67	0.90
$p = 1/40$	δ_{small}	0.71	0.94		0.67	0.89		0.74	0.95		0.67	0.85
	δ_{large}	0.67	0.90		0.67	0.89		0.73	0.94		0.65	0.79
$p = 1/260$	δ_{small}	0.72	0.94		0.67	0.83		0.74	0.95		0.68	0.87
	δ_{large}	0.69	0.91		0.67	0.81		0.74	0.94		0.66	0.82

Notes: The table shows coverage probability of A^{Bayes} and A^{MN} sets for four models subject to breaks in level (Models 1 and 2) or volatility (Models 3 and 4). Breaks occur in each time period with probability p and are of size δ_{small} or δ_{large} . The models are described in the text.

the $I(d)$ model. Moving to the results with $p > 0$, A^{Bayes} has coverage rates less than its nominal level in models 1, 2, and 4, but coverage for model 3 is similar to the results with $p = 0$; coverage rates for nominal 67% A^{MN} are approximately correct for all models, but 90% A^{MN} sets undercover by approximately 10% in models 2 and 4 when shifts are large and (in model 2) infrequent.

The conclusion we draw from these results is that the predictions sets based on asymptotic approximations developed in Sections 3 and parameterizations in Section 4 perform reasonably well in the face of the kinds of breaks that have occurred in the post-WWII U.S. macroeconomy.

5.2 Pseudo-out-of-sample Forecasts

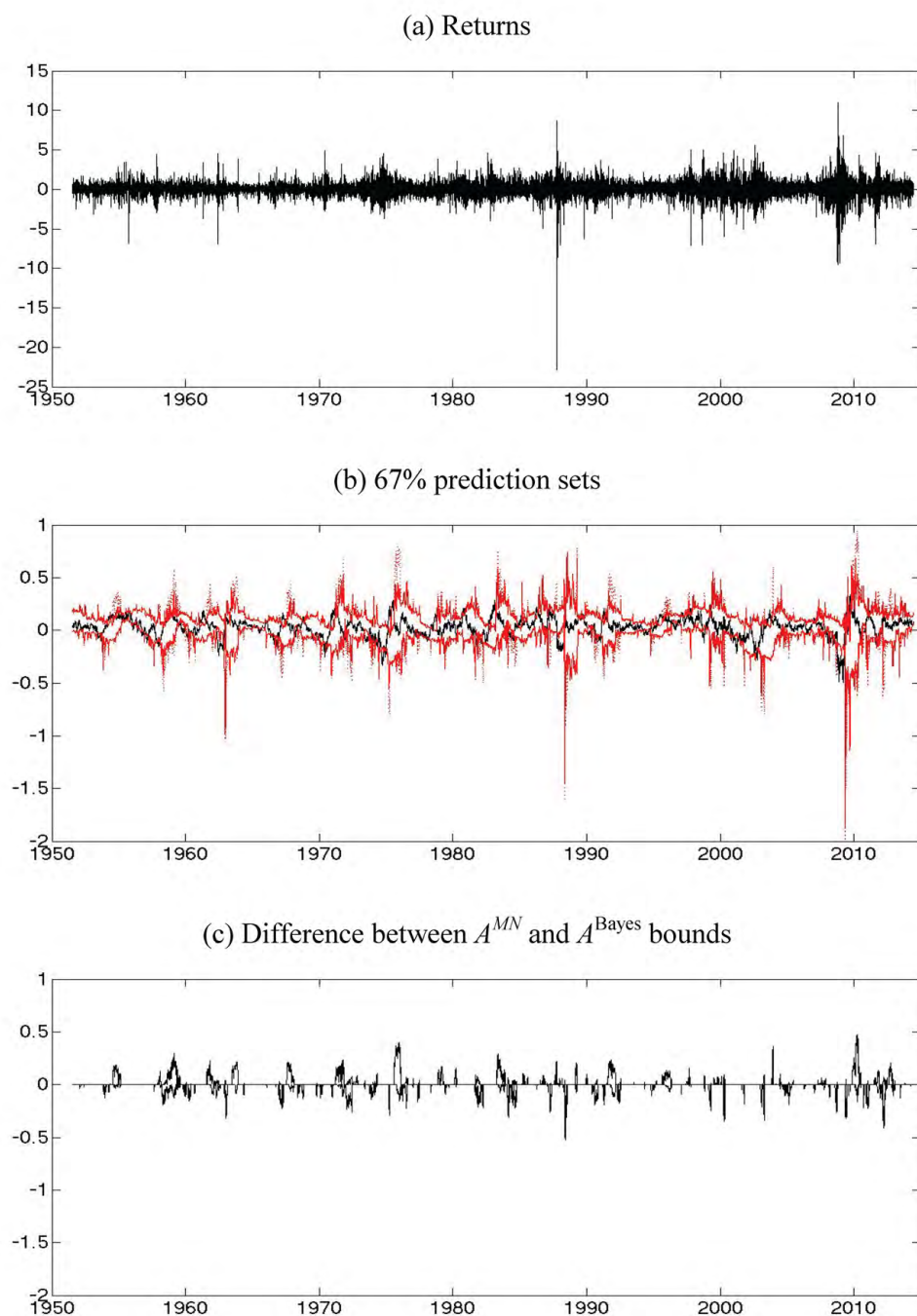
The last section examined the performance of long-run prediction sets using simulated data, but how well do the sets perform for actual data? Ideally, pseudo-out-of-sample experiments could be used to answer this question using economic time series from a wide array of stochastic processes. However, this is difficult in our setting – where we are interested in long-horizon forecasts for macroeconomic series in developed economies like the U.S. – because the available macroeconomic data provide little pseudo-out-of-sample information.

But recall that the salient definition of a long-run forecast is that the horizon is long relative to the sample data. And in contrast to macroeconomic data, there are long time series on high-frequency asset returns. One empirical test of the methods developed here is thus to see whether forecasts constructed from, say, one year of financial data, have reasonable empirical coverage for forecasts of the average value over the following half year.

Specifically, we use daily returns constructed as $r_t = 100 \ln(P_t/P_{t-1})$, where P_t is the closing value of the SP500 index; we also construct prediction sets for squared returns, r_t^2 (and thus forecast a measure of average realized volatility). We use daily data from January 3, 1950 through June 20, 2014, for a total of 16,220 returns. The pseudo-out-of-sample exercise uses a rolling sample of $T = 260$ observations to construct prediction sets for the average value of r_t and r_t^2 over the next $h = 0.5T = 130$ periods, where the choice of T matches the sample size used in the last section and in much of the empirical analysis in Section 6. Rolling through the sample in this way allows us to compute 15,831 pseudo-out-of-sample prediction sets (which, of course, are correlated because of their partial overlap).

Panel (a) of Figure 6 plots r_t over the full sample period. The series exhibits behavior typical of asset returns: little apparent short-run predictability in the level, but predictable time variation in volatility, and occasional large outliers. Panel (b) shows the 67% level A^{Bayes} and A^{MN} prediction sets along with the realized values of average returns over the

**Figure 6: Daily SP500 returns and rolling pseudo-out-of sample prediction sets
($T = 260$ and $r = 0.5$)**



Notes: Panel (a) shows daily returns for the SP500, panel (b) shows A^{Bayes} (thin red lines) and A^{MN} (red dots) prediction sets and the realization of average returns (thick blue line), and panel (c) shows the difference between the upper bounds of A^{MN} and A^{Bayes} sets (positive values) and lower bounds (negative values)

Table 6: Coverage rates for prediction sets in pseudo-out-of-sample experiment
Rolling sample, $T = 260$, $q=12$ and $r = 0.5$

Prediction Set	Returns			Squared returns	
	67%	90%		67%	90%
A^{Bayes}	0.70	0.92		0.65	0.84
A^{MN}	0.73	0.93		0.70	0.89

Notes: The table shows empirical coverage rates for A^{Bayes} and A^{MN} prediction sets in 15,831 pseudo-out-of-sample periods for the average value of SP500 returns and squared returns.

rolling forecast periods. An interesting feature of the prediction sets is how often the Bayes and frequentist sets coincide: $A^{MN} = A^{\text{Bayes}}$ approximately 70% of the time, as can also be seen from panel (c). Pseudo-out-of-sample coverage rates for 67% and 90% prediction sets are summarized in Table 6. Both A^{Bayes} and A^{MN} have sample coverage rates slightly larger than their nominal values; this result is not unexpected given the results in the preceding sections.

Squared returns are significantly more persistent than the level of returns, and are often given as an example of an economic time series that exhibits $I(d)$ low-frequency behavior (see, for instance, Ding, Granger, and Engle (1993)). Indeed, full-sample Geweke and Porter-Hudak (1983) regressions using the squared SP500 returns yield estimates of d of approximately 0.4.¹³ Table 6 indicates that the pseudo-out-of-sample coverage for A^{Bayes} is slightly lower than its nominal level, while the coverage of A^{MN} remains near its nominal level; again, these results are not unexpected given the simulation results of the last subsection, and the results in Table 1.

Of course, this exercise provides only limited information about the performance of our methods for macroeconomic variables because asset returns have substantially different properties (large outliers, relatively weak dependence in the mean, etc.). At the same time, it is encouraging to see that our method – which is in no way adapted to the specifics of these series – performs so reliably here.

6 Prediction Sets for U.S. Macroeconomic Time Series

In this section we present prediction sets for eight U.S. economic time series for forecast horizons ranging from 10 to 75 years. These series include the growth rate of per-capita values

¹³Results depend somewhat on the number of periodogram ordinate used in the regression and the treatment of outliers. For example, using $n^{1/2}$ ordinates yields $\hat{d} = 0.40$, and the same regression, but dropping the outlier associated with 1987's Black Monday yields $\hat{d} = 0.44$.

of real GDP and CPI inflation used as the running examples, and also the growth rates of real per-capita consumption expenditures, population, productivity (both total factor and labor productivity), real stock returns, and prices as measured by the PCE deflator. We construct prediction sets using post-WWII quarterly samples, and for several series, samples that extend into the early 20th century. We also examine prediction sets for inflation in Japan as a contrast to results for U.S. inflation. Sources and details of construction of the data are presented in the Data Appendix. Appendix Figures A.1-A.14 provide a variety of summary statistics for each series including a plot of the series, its low-frequency components, normalized cosine transformations, low-frequency $I(d)$ log-likelihood values, and 67% and 90% Bayes, MN and $I(0)$ prediction sets for all horizons between 10 and 75 years. Table 7 reports a summary of the prediction sets for prediction sets for 10, 25, 50, and 75 year horizons.

We now discuss the results for specific series in more detail.

Real per capita GDP. The Bayes prediction sets for per-capita GDP narrow as the forecast horizon increases, consistent with the reduction in variance of the sample mean for an $I(0)$ process. The frequentist set coincides with the Bayes set for 67% coverage and for (relatively) short horizons for 90% coverage. However, at longer horizons the 90% frequentist sets differ from Bayes sets and include smaller values of average GDP growth rates. Apparently, to guarantee high coverage uniformly in the *bcd*-model at long horizons, the frequentist sets allow for the possibility of more persistence in the GDP process, so that the slow-growth rates of the past decade are predicted to potentially persist into the future. A comparison of the prediction sets constructed using the post-WWII data and the long-annual (1901-2011) series shows that the pre-WWII data tend to widen the predictions sets, presumably because of the higher (long-run) variance in the pre-WWII data evident in Figure A.10.

At the 75-year horizon the 80% Bayes prediction interval (not shown) is 1.4 to 2.5, which roughly coincides with the 80% interval reported by the Congressional Budget Office (2005) for 75-year forecasts beginning in 2004. The coincidence of the Bayes/CBO sets arises despite important differences in the way they are computed. The CBO interval is based on simulations computed from its long-run model with inputs such as TFP growth simulated from estimated $I(0)$ models. The CBO interval differs from the Bayes interval in two important respects. First, because the simulations are carried out using fixed values of the model parameters, the CBO method ignores the parameter uncertainty in $\bar{x}_{1:T}$ (as an estimate of μ) and s_{LR}^2 (as an estimate of the long-run variance). Ignoring this uncertainty leads the CBO interval to underestimate uncertainty in the predictions. Second, in the CBO model, GDP growth is $I(0)$, while the Bayes method allows values of d that differ from $d = 0$. The log-likelihood values plotted in Figure 3 suggest that GDP growth is plausibly

Table 7: Prediction sets

a. 67% coverage

Series	Forecast horizon (in years)			
	10	25	50	75
<i>Quarterly Post-WWII Series</i>				
GDP/Pop	(1.3, 3.2)	(1.5, 2.7)	(1.6, 2.4)	(1.6, 2.4)
Cons/Pop	(1.3, 3.2)	(1.5, 2.8)	(1.6, 2.6)	(1.7, 2.5)
TF Prod	(0.2, 1.9) [0.0, 1.9]	(0.4, 1.8) [-0.3, 2.1]	(0.5, 1.8) [-0.6, 2.4]	(0.6, 1.8) [-0.8, 2.6]
Labor Prod	(1.3, 2.8)	(1.4, 2.7) [1.2, 2.8]	(1.5, 2.7) [1.0, 3.0]	(1.6, 2.7) [0.8, 3.2]
Population	(0.6, 1.1)	(0.5, 1.1) [0.4, 1.2]	(0.5, 1.2) [0.2, 1.4]	(0.4, 1.3) [0.1, 1.5]
Inflation (PCE)	(0.6, 4.1)	(0.4, 4.4)	(0.2, 4.7)	(0.0, 5.0)
Inflation (CPI)	(0.7, 4.9)	(0.7, 5.1)	(0.7, 5.3) [0.6, 5.3]	(0.6, 5.4) [0.1, 5.5]
Infl. (CPI,Japan)	(-2.8, 3.4) [-3.3, 3.6]	(-2.9, 4.2) [-4.2, 4.3]	(-3.0, 4.9) [-5.7, 5.9]	(-3.2, 5.3) [-6.9, 7.1]
Stock Returns	(0.0, 12.0)	(1.2, 11.0)	(1.8, 10.6) [-0.7, 10.7]	(2.1, 10.5) [-2.1, 12.0]
<i>Longer Span Data Series</i>				
GDP/Pop	(0.2, 4.5)	(0.9, 3.4)	(1.2, 2.9)	(1.4, 2.7)
Cons/Pop	(0.2, 2.7) [0.2, 2.9]	(0.6, 2.5) [0.6, 2.7]	(0.8, 2.4) [0.5, 3.0]	(0.9, 2.4) [0.3, 3.2]
Population	(0.6, 1.3)	(0.6, 1.4)	(0.6, 1.4)	(0.5, 1.5)
Inflation (CPI)	(-0.1, 6.1)	(0.3, 5.8)	(0.6, 5.7)	(0.6, 5.7)
Stock Returns	(1.1, 11.9)	(2.7, 10.1)	(3.4, 9.3)	(3.7, 9.0)

**Table 7: Prediction sets
(continued)**

b. 90% coverage

Series	Forecast horizon (in years)			
	10	25	50	75
<i>Quarterly Post-WWII Series</i>				
GDP/Pop	(0.4, 3.9) [0.5, 3.9]	(0.8, 3.0) [0.5, 3.0]	(1.0, 2.7) [0.0, 2.7]	(1.1, 2.6) [-0.2, 2.6]
Cons/Pop	(0.3, 3.9) [0.2, 3.9]	(0.7, 3.1) [0.2, 3.1]	(0.9, 2.9) [0.0, 2.9]	(1.0, 2.8) [-0.2, 2.8]
TF Prod	(-0.4, 2.4) [-0.8, 2.6]	(-0.3, 2.4) [-1.3, 2.9]	(-0.3, 2.4) [-1.8, 3.4]	(-0.3, 2.5) [-2.2, 3.8]
Labor Prod	(0.6, 3.4) [0.2, 3.6]	(0.7, 3.3) [0.2, 3.6]	(0.7, 3.3) [-0.2, 4.0]	(0.7, 3.4) [-0.5, 4.4]
Population	(0.4, 1.2) [0.4, 1.3]	(0.3, 1.4) [0.1, 1.5]	(0.2, 1.5) [-0.2, 1.8]	(0.0, 1.6) [-0.4, 2.0]
Inflation (PCE)	(-0.8, 5.4) [-1.0, 5.4]	(-1.4, 6.0) [-2.3, 6.4]	(-2.2, 6.8) [-4.1, 8.2]	(-2.9, 7.4) [-5.4, 9.5]
Inflation (CPI)	(-1.0, 6.4) [-1.2, 6.4]	(-1.4, 6.8) [-2.2, 7.1]	(-2.0, 7.5) [-3.9, 8.8]	(-2.4, 8.0) [-5.2, 10.1]
Infl. (CPI, Japan)	(-5.4, 5.6) [-6.6, 6.0]	(-6.3, 6.7) [-8.9, 8.2]	(-7.5, 8.1) [-12.1, 11.3]	(-8.4, 9.1) [-14.8, 14.0]
Stock Returns	(-5.0, 16.6) [-5.1, 16.6]	(-3.9, 15.3) [-7.6, 16.2]	(-3.7, 15.2) [-10.5, 19.3]	(-3.7, 15.4) [-13.1, 21.9]
<i>Longer Span Data Series</i>				
GDP/Pop	(-1.4, 6.1) [-1.6, 6.5]	(-0.1, 4.3) [0.1, 4.3]	(0.4, 3.6) [0.1, 3.6]	(0.6, 3.3) [0.1, 3.3]
Cons/Pop	(-0.8, 3.6) [-0.8, 4.3]	(-0.3, 3.2) [-0.6, 3.6]	(-0.1, 3.0) [-1.0, 4.0]	(-0.0, 3.0) [-1.3, 4.3]
Population	(0.3, 1.6) [0.1, 1.6]	(0.2, 1.6) [0.1, 1.6]	(0.1, 1.7) [-0.1, 1.9]	(0.0, 1.8) [-0.3, 2.0]
Inflation (CPI)	(-2.5, 8.5) [-3.2, 8.8]	(-2.2, 8.2) [-3.2, 8.8]	(-2.5, 8.6) [-3.2, 8.8]	(-2.9, 9.0) [-3.9, 9.7]
Stock Returns	(-3.1, 16.1) [-3.1, 16.1]	(-0.8, 13.4) [-0.8, 13.4]	(0.1, 12.5) [-2.0, 12.5]	(0.4, 12.2) [-3.5, 13.5]

Notes: This table shows the 67% and 90% prediction sets for forecast horizons, $h = 10, 25, 50$, and 75 years. The A^{Bayes} sets are shown in parentheses and are based on the $I(d)$ model with uniform prior for $-0.4 \leq d \leq 1.0$. The A^{MN} sets are shown in brackets, and are omitted if they coincide with the A^{Bayes} sets. By construction the A^{MN} sets control asymptotic coverage in the bcd -model with $-0.4 \leq d \leq 1.0$, and b and c unrestricted.

characterized by a process with some low-frequency anti-persistence, and this translates into less forecast uncertainty than the CBO's $I(0)$ model. Thus, the CBO method tends to understate forecast uncertainty because it ignores parameter uncertainty in the estimated mean and long-run variance, and to overstate forecast uncertainty because its model does not capture long-run anti-persistence associated with negative values of d . Apparently these two errors cancel, so that the CBO prediction interval essentially coincides with the Bayes set.

Productivity. The log-likelihood values for d indicate that productivity (TFP and average labor productivity) may have somewhat greater than $I(0)$ persistence; see appendix Figures A.3 and A.4. This translates into prediction sets that are wider than $I(0)$ sets, particularly for frequentist sets at large forecast horizons. Bayes intervals are essentially flat as the forecast horizon increases (unlike in an $I(0)$ model, where the intervals narrow), while the frequentist sets widen (the unmodified Bayes intervals systematically undercover for larger values of d , forcing the frequentist intervals to more heavily weigh the possibility of larger d).

Population. U.S. population growth shows considerable low-frequency variability over the 20th century and the post-WWII period.¹⁴ Immigration and fertility dynamics are presumably at the source of these long swings. The low-frequency MLE of d is very close to unity over both sample periods, with the $I(1)$ log-likelihood more than 7 points higher than in the $I(0)$ model. Table 7 and appendix Figures A.5 and A.12 show prediction intervals that widen as the forecast horizon increases, a familiar characteristic of $I(1)$ predictive densities. There is little difference in the sets constructed using the post-WWII samples and long-samples.

Inflation. As discussed above, the inflation process is characterized by more than $I(0)$ persistence, and this is reflected in the prediction sets in two ways. First, they are not centered at the sample mean of the series, but rather at a level dictated by the values near the end of sample period, and second, the prediction sets widen with the forecast horizon. The prediction intervals indicate considerable uncertainty in inflation even at relatively short horizons; this is true for Bayes and frequentist sets (which essentially coincide). For example, while the 10-year 67% Bayes prediction set for U.S. CPI inflation is $(0.7, 4.9)$, the 90% set widens to $(-1.0, 6.4)$.

These predictions sets may strike some readers as too large, but it is instructive to consider the history of Japan where the 10-year moving average of CPI inflation was less than zero

¹⁴The quarterly post-WWII population series shown in Appendix Figure A.5 is not seasonally adjusted and shows substantial seasonal variation. Because we focus on low-frequency transformations of the series, this seasonality does not affect the empirical results.

from 2003 through the end of the sample. (See appendix Figure A.8.) Moreover, they are in line with predictive densities derived from asset prices. For example, Kitsul and Wright (2012) use CPI-based derivatives to compute market-based risk-neutral predictive densities for 10-year ahead average values of inflation. They find deflation (average inflation less than 0%) probabilities that averaged approximately 15% over 2011 and “high inflation” (average inflation greater than 4%) of 30%.¹⁵ The corresponding probabilities computed from the Bayes predictive density constructed using the post-WWII data are 11% for deflation and 28% for high inflation.

Stock Returns. Post-WWII real stock returns exhibit slightly more persistence than is implied by the $I(0)$ model, and this translates into prediction sets that are wider than implied by the $I(0)$ model. For example, at the 25-year horizon, the 67%- $I(0)$ prediction set (from appendix Figure A.9) is (3.1, 10.3) while the corresponding Bayes and MN prediction sets (from Table 7) are (1.2, 11.0). The longer-span data suggest somewhat less persistence ($\hat{d}^{MLE} = 0.0$ for the 1926-2001 sample) yielding Bayes and frequentist prediction intervals that are somewhat narrower than those constructed using the post-WWII data.

Pastor and Stambaugh (2012) survey the large literature on long-run stock return volatility and construct Bayes predictive densities using models that allow for potentially persistent components in returns and incorporate parameter uncertainty. While their results rely on more parametric models than ours—they use all frequencies and exact Gaussian likelihoods—our empirical conclusions are similar. Using our notation, Pastor and Stambaugh (2012) are concerned with the behavior of the variance of $\sqrt{h}\bar{x}_{T+1:T+h}$ and how this variance changes with the forecast horizon h . If the variance of $\sqrt{h}\bar{x}_{T+1:T+h}$ is unchanged as h increases, and if the predictive density is Gaussian, then the width of prediction intervals for $\bar{x}_{T+1:T+h}$ will be proportional to $h^{-1/2}$. Pastor and Stambaugh find that the variance of $\sqrt{h}\bar{x}_{T+1:T+h}$ is not constant, but rather increases with h . Consistent with this, we find Bayes prediction sets that narrow as h increases, but more slowly than $h^{-1/2}$.

Results for different values of q . As discussed in Sections 2 and 4, the choice of $q = 12$ involved an efficiency/robustness trade-off, where a larger value of q results in more information about the scale and shape parameter, but potential misspecification because the higher-frequency spectrum may not be well-described by the same model and parameter. It is therefore interesting to see how the prediction sets vary with q , and this is reported in Table 8, which shows the 67% and 90% prediction sets for the 25-year ahead forecasts for $q = 6, 12$, and 24. Looking across all of the entries, the prediction sets behave roughly

¹⁵See Kitsul and Wright (2012), Figures 3 and 4. Fleckenstein, Longstaff, and Lustig (2013) estimate somewhat lower probabilities for deflation, but similar probabilities for inflation exceeding 4%. (See their Figures 4 and 5). For a related calculation, see Figure 3 in Hilsher, Raviv, and Reis (2014).

Table 8: Prediction sets for various values of q , 25-year horizon

	67 %				90 %		
	$q = 6$	$q = 12$	$q = 24$		$q = 6$	$q = 12$	$q = 24$
<i>Quarterly Post-WWII Data</i>							
GDP/Pop	(0.5, 2.4)	(1.5, 2.7)	(1.5, 2.6)		(-0.4, 2.9) [-1.0, 2.9]	(0.8, 3.0) [0.5, 3.0]	(1.0, 3.0) [0.3, 3.0]
Cons/Pop	(1.0, 2.7)	(1.5, 2.8)	(1.6, 2.7)		(-0.1, 3.1) [-0.5, 3.1]	(0.7, 3.1) [0.2, 3.1]	(1.0, 3.1) [0.0, 3.1]
TF Prod	(0.2, 1.6) [-0.2, 1.8]	(0.4, 1.8) [-0.3, 2.1]	(0.4, 1.8) [-0.4, 2.2]		(-0.4, 2.1) [-1.0, 2.5]	(-0.3, 2.4) [-1.3, 2.9]	(0.1, 2.3) [-1.2, 3.0]
Labor Prod	(1.4, 2.7) [1.2, 2.7]	(1.4, 2.7) [1.2, 2.8]	(1.6, 2.7) [1.1, 2.9]		(0.7, 3.2) [0.2, 3.5]	(0.7, 3.3) [0.2, 3.6]	(1.0, 3.1) [0.3, 3.6]
Population	(0.6, 1.3) [0.4, 1.3]	(0.5, 1.1) [0.4, 1.2]	(0.4, 1.0) [0.3, 1.1]		(0.3, 1.5) [0.0, 1.7]	(0.3, 1.4) [0.1, 1.5]	(0.2, 1.2) [0.1, 1.3]
Inflation (PCE)	(0.7, 4.8)	(0.4, 4.4)	(1.1, 4.4) [0.2, 4.4]		(-1.3, 6.5) [-2.7, 6.9]	(-1.4, 6.0) [-2.3, 6.4]	(-0.4, 5.8) [-2.5, 6.8]
Inflation (CPI)	(0.8, 5.4)	(0.7, 5.1)	(1.4, 5.0) [1.0, 5.0]		(-1.5, 7.2) [-3.0, 7.6]	(-1.4, 6.8) [-2.2, 7.1]	(-0.1, 6.5) [-2.7, 7.6]
Infl. (CPI,Japan)	(-2.0, 5.0) [-5.3, 5.8]	(-2.9, 4.2) [-4.2, 4.3]	(-2.2, 4.4)		(-5.2, 7.8) [-9.7, 10.3]	(-6.3, 6.7) [-8.9, 8.2]	(-5.3, 7.1) [-7.4, 7.1]
Stock Returns	(0.2, 11.1)	(1.2, 11.0)	(2.7, 10.8) [1.8, 10.8]		(-5.7, 15.6) [-8.8, 15.6]	(-3.9, 15.3) [-7.6, 16.2]	(-0.9, 14.4) [-6.9, 16.9]
<i>Longer Span Data Series</i>							
GDP/Pop	(0.8, 2.9)	(0.9, 3.4)	(0.8, 3.2)		(-0.2, 3.8)	(-0.1, 4.3)	(-0.3, 4.1)
Cons/Pop	(0.7, 2.5)	(0.6, 2.5) [0.6, 2.7]	(0.9, 2.6) [0.9, 2.7]		(-0.1, 3.3) [-0.3, 3.4]	(-0.3, 3.2) [-0.6, 3.6]	(0.2, 3.3) [-0.2, 3.7]
Population	(0.8, 1.6)	(0.6, 1.4)	(0.5, 1.3)		(0.4, 1.9) [0.3, 2.0]	(0.2, 1.6) [0.1, 1.6]	(0.2, 1.5) [0.1, 1.5]
Inflation (CPI)	(0.3, 6.1)	(0.3, 5.8)	(0.8, 5.3)		(-2.3, 8.7)	(-2.2, 8.2)	(-1.1, 7.0) [-1.6, 7.5]
Stock Returns	(1.9, 10.8)	(2.7, 10.1)	(3.5, 9.5)		(-2.7, 14.5)	(-0.8, 13.4)	(1.2, 11.8) [0.6, 11.8]

Notes: A^{Bayes} sets are shown in parentheses and A^{MN} sets are shown in brackets when they differ from Bayes sets. See notes to Table 7 for additional information.

as expected, in the sense that they remain centered at roughly the same value but tend to narrow as q increases. For example, averaging across the 14 series, the 67% MN prediction set is 7% narrower using $q = 24$ than with $q = 12$ broadly consistent the results in Table 3. There are two noteworthy exceptions. First, per-capita GDP (quarterly post-WWII) shows a lower bound for the $q = 6$ prediction sets that is one percentage point lower than the corresponding value for the $q = 12$ sets. This arises because the $q = 6$ log-likelihood functions for d are much flatter than the $q = 12$ functions summarized in Table 3. Thus, when $q = 6$, the prediction sets account for the possibility of more persistence in the series and this includes an extrapolation of the low-growth levels experienced by the U.S. economy over the past decade. The second exception is quarterly U.S. inflation. Here the lower bounds of the $q = 24$ Bayes 90%-prediction sets are more than one percentage point higher than for $q = 12$. This arises because the $q = 24$ likelihood values suggest less persistence than $q = 12$ value. (When $q = 24$ the likelihood peaks at values of d around 0.4 while the peak is around $d = 0.7$ for $q = 12$.)

7 Conclusions

This paper has considered the problem of quantifying uncertainty about long-run predictions using prediction sets that contain the realized future value of a variable of interest with prespecified probability. The long-run nature of the problem both simplifies and complicates the problem relative to short-run predictions. The problem is simplified because of our focus on forecasting long-run averages using a relatively small number of (low-frequency) weighted averages of the sample data. We show that these averages have an approximate joint normal distribution, which greatly simplifies the prediction problem. However, the prediction problem is complicated because the covariance matrix of the limiting normal distribution depends on the spectrum over very low frequencies, and there is limited sample information about its shape. Uncertainty about the low-frequency characteristics of the stochastic process is then an important component of the uncertainty about long-run predictions.

We proposed a flexible parametric model (the *bcd*-model) to characterize the shape of the spectrum at low frequencies. Uncertainty about the shape then becomes equivalent to uncertainty about the values of the *bcd*-parameters. Incorporating this parameter uncertainty into prediction uncertainty is straightforward in a Bayesian framework, and we provide the details in the context of the long-run prediction problem. However, because of the paucity of sample information about these long-run parameters, the resulting Bayes prediction sets may depend importantly on the specifics of the prior. This motivates us to robustify the Bayes sets by enlarging them so that, by construction, they control coverage uniformly over all values

of the bcd -parameters. We construct minimum expected length frequentist prediction sets using an approximate “least favorable distribution” for the parameters, and we generalize these to conditionally sensible frequentist prediction sets using ideas from Müller and Norets (2012).

We apply these methods and construct prediction sets for nine macroeconomic time series for forecast horizons of up to 75 years. In general, we found that for many series, the prediction sets are wider than those that one obtains from the $I(0)$ model, but narrower than one would obtain from, say, the $I(1)$ model. From a statistical point of view, this underlines the importance of modelling the spectral shape at low frequencies in a flexible manner. Substantively, it demonstrates that even after accounting for a wide variety of potential long-run instabilities and dependencies, it is still possible to make informative probability statements about (very) long-run forecasts.

Our analysis has been univariate in the sense that we have constructed predictions sets for a scalar random variable $\bar{x}_{T+1:T+h}$ using sample values of x_t . However, answers to some questions require multivariate prediction sets. The statistical theory discussed and developed in Section 3 carries over directly to multivariate settings. That said, there are important practical obstacles to constructing multivariate prediction sets. These obstacles include finding a convenient, but flexible, parameterization of the multivariate local-to-zero spectrum, constructing accurate approximations to least favorable distributions with high dimensional θ , and computing accurate approximations to the density of relevant invariants. Overcoming these obstacles is left to future research.

References

- BARRO, R. J. (2006): “Rare disasters and asset markets in the twentieth century,” *The Quarterly Journal of Economics*, 121(3), 823–866.
- BERAN, J. (1994): *Statistics for Long-Memory Processes*. Chapman and Hall, London.
- CAMPBELL, J. Y., AND L. M. VICEIRA (1999): “Consumption and Portfolio Decisions When Expected Returns are Time Varying,” *Quarterly Journal of Economics*, 114, 433–495.
- CARTER, S. B., S. S. GARTNER, M. R. HAINES, A. L. OLMSTEAD, R. SUTCH, AND G. WRIGHT (2006): *Historical Statistics of the United States, Millennial Edition Online*. Cambridge University Press.

- COGLEY, T., AND T. SARGENT (2014): “Measuring Price Level Uncertainty and Stability in the U.S., 1850-2012,” *forthcoming in the Review of Economics and Statistics*.
- CONGRESSIONAL BUDGET OFFICE (2005): “Quantifying Uncertainty in the Analysis of Long-Term Social Security Projections,” *CBO Background paper*.
- DAVIDSON, J. (1994): *Stochastic Limit Theory*. Oxford University Press, New York.
- DING, Z., C. W. J. GRANGER, AND R. F. ENGLE (1993): “A Long Memory Property of Stock Market Returns and a New Model,” *Journal of Empirical Finance*, 1, 83–116.
- DOORNIK, J. A., AND M. OOMS (2004): “Inference and Forecasting for ARFIMA Models With an Application to US and UK Inflation,” *Studies in Nonlinear Dynamics & Econometrics*, 8.
- ELLIOTT, G. (2006): “Forecasting with Trending Data,” in *Handbook of Economic Forecasting, Volume 1*, ed. by C. W. J. G. G. Elliott, and A. Timmerman. North-Holland.
- ELLIOTT, G., U. K. MÜLLER, AND M. W. WATSON (2014): “Nearly Optimal Tests When a Nuisance Parameter is Present Under the Null Hypothesis,” *forthcoming in Econometrica*.
- FERNALD, J. (2012): “A Quarterly, Utilization-Adjusted Series on Total Factor Productivity,” *FRBSF Working Paper 2012-9*.
- FLECKENSTEIN, M., F. LONGSTAFF, AND H. LUSTIG (2013): “Deflation Risk,” *NBER working paper w19238*.
- GEWEKE, J., AND S. PORTER-HUDAK (1983): “The Estimation and Application of Long Memory Time Series Models,” *Journal of Time Series Analysis*, 4, 221–238.
- GRANGER, C. W., AND Y. JEON (2007): “Long-term forecasting and evaluation,” *International Journal of Forecasting*, 23(4), 539 – 551.
- HILSHER, J., A. RAVIV, AND R. REIS (2014): “Inflation Away the Public Debt? An Empirical Assessment,” *NBER Working Paper w 20339*.
- KEMP, G. C. R. (1999): “The Behavior of Forecast Errors from a Nearly Integrated AR(1) Model as Both Sample Size and Forecast Horizon Become Large,” *Econometric Theory*, 15(2), pp. 238–256.
- KIEFER, N., AND T. J. VOGELSANG (2002): “Heteroskedasticity-Autocorrelation Robust Testing Using Bandwidth Equal to Sample Size,” *Econometric Theory*, 18, 1350–1366.

- (2005): “A New Asymptotic Theory for Heteroskedasticity-Autocorrelation Robust Tests,” *Econometric Theory*, 21, 1130–1164.
- KIEFER, N. M., T. J. VOGELSANG, AND H. BUNZEL (2000): “Simple Robust Testing of Regression Hypotheses,” *Econometrica*, 68, 695–714.
- KITSUL, Y., AND J. H. WRIGHT (2012): “The Economics of Options-Implied Inflation Probability Density Functions,” *Working Paper, Johns Hopkins University*.
- LEE, R. (2011): “The Outlook for Population Growth,” *Science*, 333, 569–573.
- LINDGREN, G. (2013): *Stationary Stochastic Processes: Theory and Applications*. Chapman & Hall/CRC.
- LUKE, Y. (1975): *Mathematical Functions and Their Approximations*. Academic Press, New York.
- MUIRHEAD, R. J. (1982): *Aspects of Multivariate Statistical Theory*. Wiley.
- MÜLLER, U. K. (2014): “HAC Corrections for Strongly Autocorrelated Time Series,” *Journal of Business and Economic Statistics*, 32, 311–322.
- MÜLLER, U. K., AND A. NORETS (2012): “Credibility of Confidence Sets in Nonstandard Econometric Problems,” *Working Paper, Princeton University*.
- MÜLLER, U. K., AND M. W. WATSON (2008): “Testing Models of Low-Frequency Variability,” *Econometrica*, 76, 979–1016.
- (2013): “Low-Frequency Robust Cointegration Testing,” *Journal of Econometrics*, 174, 66–81.
- PASTOR, L., AND R. F. STAMBAUGH (2012): “Are Stocks Really Less Volatile in the Long Run?,” *Journal of Finance*, LXVII, 431–477.
- PESAVENTO, E., AND B. ROSSI (2006): “Small-Sample Confidence Intervals for Multivariate Impulse Response Functions at Long Horizons,” *Journal of Applied Econometrics*, 21(8), 1135–1155.
- PHILLIPS, P. C. B. (1998): “Impulse Response and Forecast Error Variance Asymptotics in Nonstationary VARs,” *Journal of Econometrics*, 83, 21–56.

- RAFTERY, A. E., N. LI, H. SEVČÍKOVÁ, P. GERLAND, AND G. K. HEILIG (2012): “Bayesian probabilistic population projections for all countries,” *Proceedings of the National Academy of Sciences*.
- RIETZ, T. A. (1988): “The equity risk premium a solution,” *Journal of monetary Economics*, 22(1), 117–131.
- SIEGEL, J. (2007): *Stocks for the Long Run: The Definitive Guide to Financial Market Returns and Long Term Investment Strategies*. McGraw-Hill.
- STOCK, J. H. (1996): “VAR, Error Correction and Pretest Forecasts at Long Horizons,” *Oxford*, 58, 685–701.
- (1997): “Cointegration, Long-Run Comovements, and Long-Horizon Forecasting,” in *Advances in Econometrics: Proceedings of the Seventh World Congress of the Econometric Society*, ed. by D. Kreps, and K. Wallis, pp. 34–60, Cambridge. Cambridge University Press.
- STOCK, J. H., AND M. W. WATSON (2002): “Has the Business Cycle Changed and Why?,” in *NBER Macroeconomics Annual 2002*, ed. by M. Gertler, and K. S. Rogoff, pp. 159–218. MIT Press, Cambridge, MA.
- (2007): “Why Has Inflation Become Harder to Forecast?,” *Journal of Money, Credit, and Banking*, 39, 3–34.
- WATSON, M. W. (2014): “Inflation Persistence, the NAIRU, and the Great Recession,” *American Economic Review*, 104, 31–36.

Data Appendix

Series	Sample period	Description	Sources and Notes
<i>Post- WWII Quarterly Series</i>			
GDP	1947:Q1-2012:Q1	Real GDP (Billions of Chained 2005 Dollars)	FRED Series: GDPC96
Consumption	1947:Q1-2012:Q1	Real Personal Consumption Expenditures (Billions of Chained 2005 Dollars)	FRED Series PCECC96
Total Factor Productivity	1947:Q2-2012:Q1	Growth Rate of TFP	From John Fernald's Web Page: Filename Quarterly)TFP.XLSX, series name DTFP. The series is described in Fernald (2012)
Labor Productivity	1947:Q1-2012:Q1	Output per hour in the non-farm business sector	FRED Series: OPHNFB
Population	1947:Q1-2012:Q1	Population (mid Quarter) (Thousands)	Bureau of Economic Analysis NIPA Table 7.1. Adjusted for an outlier in 1960:Q1
Prices: PCE Deflator	1947:Q1-2012:Q1	PCE deflator	FRED Series: PCECTPI
Inflation (CPI)	1947:Q1-2012:Q1	CPI	FRED Series: CPIAUCSL. The quarterly price index was computed as the average of the monthly values
Inflation (CPI, Japan)	1960:Q1-2012:Q1	CPI for Japan	FRED Series: CPI_JPN. The quarterly price index was computed as the average of the monthly values
Stock Returns	1947:Q1-2012:Q1	CRSP Real Value-Weighted Returns	CRSP Nominal Monthly Returns are from WRDS. Monthly real returns were computed by subtracting the change in the logarithm in the CPI from the nominal returns, which were then compounded to yield quarterly returns. Values shown are 400×the logarithm of gross quarterly real returns.
<i>Longer Span Data Series</i>			
Real GDP	1900-2011	Real GDP (Billions of Chained 2005 Dollars)	<u>1900-1929</u> : Carter, Gartner, Haines, Olmstead, Sutch, and Wright (2006), Table Ca9: Real GDP in \$1996 <u>1929-2011</u> : BEA Real GDP in \$2005. Data were linked in 1929
Real Consumption	1900-2011	Real Personal Consumption Expenditures (Billions of Chained 2005 Dollars)	<u>1900-1929</u> : Carter, Gartner, Haines, Olmstead, Sutch, and Wright (2006), Table Cd78 Consumption expenditures, by type, Total, \$1987 <u>1929-2011</u> : BEA Real Consumption in \$2005. Data were linked in 1929
Population	1900-2011	Population	<u>1900-1949</u> : Carter, Gartner, Haines, Olmstead, Sutch, and Wright (2006), Table Aa6 (Population Total including Armed Forces Oversees), 1917-1919 and 1930-1959; Table Aa7 (Total, Resident) 1900-1916 and 1920-1929) <u>1950-2011</u> : Bureau of the Census (Total, Total including Armed Forces Oversees),
Inflation (CPI)	1913-2011	Consumer Price Index	BLS Series: CUUR0000SA0
Stock Returns	1926:Q1-2012:Q1	CRSP Real Value-Weighted Returns	Described above

Generally, the data used in the paper are growth rates computed as differences in logarithms in percentage points at annual rate ($400 \times \ln(X_t/X_{t-1})$ if X is measured quarterly and $100 \times \ln(X_t/X_{t-1})$ if X is measured annually). The exceptions are stock returns which are $400 \times X_t$, where X_t is the gross quarterly real return (that is, $X_t = 1 + R_t$, where R_t is the net return), and TFP which is reported in percentage points at annual rate in the source data.

8 Appendix

8.1 Central Limit Theorem of Section 3

Theorem 1 Let $\Delta x_{T,t} = \sum_{s=-\infty}^{\infty} c_{T,s} \varepsilon_{t-s}$. Suppose that

(i) $\{\varepsilon_t, \mathcal{F}_t\}$ is a martingale difference sequence with $E(\varepsilon_t^2) = 1$, $\sup_t E(|\varepsilon_t|^{2+\delta}) < \infty$ for some $\delta > 0$, and

$$E(\varepsilon_t^2 - 1 | \mathcal{F}_{t-m}) \leq \xi_m \quad (21)$$

for some sequence $\xi_m \rightarrow 0$.

(ii) For every $\epsilon > 0$ there exists an integer $L_\epsilon > 0$ such that $\limsup_{T \rightarrow \infty} T^{-1} \sum_{l=L_\epsilon T+1}^{\infty} \left(T \sup_{|s| \geq l} |c_{T,s}| \right)^2 < \epsilon$.

(iii) $\sum_{s=-\infty}^{\infty} c_{T,s}^2 < \infty$ (but not necessarily uniformly in T). The spectral density of $\Delta x_{T,t}$ thus exists; denote it by $F_T : [-\pi, \pi] \mapsto \mathbb{R}$.

(iii.a) Assume that there exists a function $S : \mathbb{R} \mapsto \mathbb{R}$ such that $\omega \mapsto \omega^2 S(\omega)$ is integrable, and for all fixed K ,

$$\int_0^K |F_T(\frac{\omega}{T}) - \omega^2 S(\omega)| d\omega \rightarrow 0. \quad (22)$$

(iii.b) For every diverging sequence $K_T \rightarrow \infty$

$$T^{-3} \int_{K_T/T}^{\pi} F_T(\lambda) \lambda^{-4} d\lambda = \int_{K_T}^{\pi T} F_T(\omega/T) \omega^{-4} d\omega \rightarrow 0. \quad (23)$$

(iii.c)

$$T^{-3/2} \int_{1/T}^{\pi} F_T(\lambda)^{1/2} \lambda^{-2} d\lambda = T^{-1/2} \int_1^{\pi T} F_T(\omega/T)^{1/2} \omega^{-2} d\omega \rightarrow 0. \quad (24)$$

(iv) For some fixed integer H , the bounded and integrable function $g : [0, H] \mapsto \mathbb{R}$ satisfies $\int_0^H g(s) ds = 0$, and with $G(r) = \int_0^r g(s) ds$, for some constant C , $|\sum_{t=1}^{HT} e^{-i\lambda t} G(\frac{t-1}{T})| \leq C \lambda^{-2} T^{-1}$ uniformly in T and λ .

Then

$$T^{-1/2} \int_0^H g(s) x_{T,[sT]+1} ds \Rightarrow \mathcal{N}(0, \int_{-\infty}^{\infty} S(\omega) \left| \int_0^H e^{-i\omega s} g(s) ds \right|^2 d\omega) \quad (25)$$

where $x_{T,t} = \sum_{s=1}^t \Delta x_{T,s}$.

Remarks: Note that the linear process $\Delta x_{T,t}$ is not restricted to be causal. The m.d.s. structure of the driving errors ε_t in assumption (i) allows for some departures from strict stationarity. It also accommodates conditional heteroskedasticity, with the second order dependence limited by the mixingale condition (21).

With the pseudo-spectrum of $x_{T,t}$ defined as $R_T(\lambda) = F_T(\lambda)/|1 - e^{-i\lambda}|^2$, assumption (iii.a) is equivalent to

$$T^2 \int_0^K \omega^2 |T^{-2} R_T(\frac{\omega}{T}) - S(\omega)| d\omega \rightarrow 0 \quad (26)$$

since for any fixed K , $\sup_{0 \leq \omega \leq K} |T^{-2} \frac{\omega^2}{|1 - e^{-i\omega/T}|^2} - 1| \rightarrow 0$. Thus, (iii.a) is equivalent to the convergence of the pseudo-spectrum of $x_{T,t}$ to S in a T^{-1} neighborhood of the origin in the sense of (26).

To better understand the role of assumptions (ii) and (iii), consider some leading examples. Suppose first that $\Delta x_{T,t}$ is causal and weakly dependent with exponentially decaying $c_{T,s}$, $|c_{T,s}| \leq C_0 e^{-C_1 s}$ for some $C_0, C_1 > 0$, as would arise in causal and invertible ARMA models of any fixed and finite order. Then $T^{-1} \sum_{l=LT+1}^{\infty} \left(T \sup_{|s| \geq l} |c_{T,s}| \right)^2 \rightarrow 0$ for any $L > 0$, $\omega^2 S(\omega)$ is constant and equal to the long-run variance of $\Delta x_{T,t}$, and (23) and (24) hold, since F_T is bounded, $\int_{K_T}^{\infty} \omega^{-4} d\omega \rightarrow 0$ for any $K_T \rightarrow \infty$ and $\int_1^{\infty} \omega^{-2} d\omega < \infty$.

Second, suppose $\Delta x_{T,t}$ is fractionally integrated with parameter $d \in (-1/2, 1/2)$ (corresponding to $x_{T,t}$ being fractionally integrated of order $d+1$). With $\Delta x_{T,t}$ scaled by T^{-d} , $c_{T,s} \approx C_0 T^{-d} s^{d-1}$, so that $T^{-1} \sum_{l=LT+1}^{\infty} \left(T \sup_{|s| \geq l} |c_{T,s}| \right)^2 \rightarrow \int_L^{\infty} s^{2d-2} ds$, which can be made arbitrarily small by choosing L large. Further, for λ close to zero, $F_T(\lambda) \approx C_0^2 (\lambda T)^{-2d}$, so that $\omega^2 S(\omega) = C_0^2 \omega^{-2d}$, and (23) and (24) are seen to hold under weak assumptions about higher frequency properties of $\Delta x_{T,t}$. For instance, even integrable poles in F_T at frequencies other than zero can be accommodated.

Third, suppose $x_{T,t}$ is an AR(1) process with local-to-unity coefficient $\rho_T = 1 - c/T$ and unit innovation variance. Then $c_{T,0} = 1$ and $c_{T,s} = -(1 - \rho_T) \rho_T^s$, $s > 0$. Thus $T^{-1} \sum_{l=LT+1}^{\infty} \left(T \sup_{|s| \geq l} |c_{T,s}| \right)^2 \rightarrow c^2 \int_L^{\infty} e^{-2cs} ds$, which can be made arbitrarily small by choosing L large. Further, $F_T(\lambda) = |1 - e^{-i\lambda}|^2 / |1 - \rho_T e^{-i\lambda}|^2$, which is seen to satisfy (22) with $S(\omega) = 1/(\omega^2 + c^2)$. Conditions (23) and (24) also hold in this example, since $F_T(\lambda) \leq 1$.

As a final example, suppose $\Delta x_{T,t} = T\varepsilon_t - T\varepsilon_{t-1}$ (inducing $x_{T,t}$ to be i.i.d. conditional on ε_0). Here $F_T(\lambda) = T^2 |1 - e^{-i\lambda}|^2 = 4T^2 \sin(\lambda/2)^2$, so that $S(\omega) = 1$, and (23) evaluates to $4 \int_{K_T}^{\pi T} T^2 \sin(\omega/2T)^2 \omega^{-4} d\omega \leq \int_{K_T}^{\pi T} \omega^{-2} d\omega \rightarrow 0$, and (23) to $2T^{-1/2} \int_1^{\pi T} T \sin(\frac{1}{2}\omega/T) \omega^{-2} d\omega \leq T^{-1/2} \int_1^{\pi T} \omega^{-1} d\omega \rightarrow 0$, where the inequalities follow from $\sin(\lambda) \leq \lambda$ for all $\lambda \geq 0$.

Assumption (iv) of Theorem 1 specifies the sense in which the weights g in (25) are smooth, so that the properties of the $g(s)$ -weighted average of $x_{T,[sT]+1}$ are wholly determined by the low-frequency properties of $x_{T,t}$. The number H is assumed to be an integer to ease notation. Note that a constant g would not satisfy assumption (iv), as it does not integrate to zero, but all functions of interest in the context of this paper do.

Lemma 1 *For some $0 < r < H-1$, let $g_{q+1} : [0, H] \mapsto \mathbb{R}$ equal $g_{q+1}(s) = -\mathbf{1}[0 \leq s \leq 1] + r^{-1} \mathbf{1}[1 < s \leq 1+r]$ and let $g_j : [0, H] \mapsto \mathbb{R}$ equal to $g_j(s) = \mathbf{1}[s \leq 1] \sqrt{2} \cos(\pi j s)$ for $j = 1, \dots, q$. Then for any fixed and real λ_j , $g(s) = \sum_{j=1}^{q+1} \lambda_j g_j(s)$ satisfies assumption (iv) of Theorem 1.*

Proof. It is clearly sufficient to show that each g_j satisfies the assumption. This follows from a direct computation. ■

The implication of Theorem 1 that is of interest for Section 3 follows from the following Corollary.

Corollary 1 Suppose $g_1, \dots, g_{q+1}$ are as in Lemma 1. Then under the assumptions of Theorem 1 (i)-(iii),

$$T^{-1/2} \int_0^H \begin{bmatrix} g_1(s) \\ \vdots \\ g_q(s) \\ g_{q+1}(s) \end{bmatrix} x_{T, \lfloor sT \rfloor + 1} ds \Rightarrow \mathcal{N}(0, \Sigma)$$

where $\Sigma_{j,k} = \int_{-\infty}^{\infty} S(\omega) \left(\int_0^H e^{-i\omega s} g_j(s) ds \right) \left(\int_0^H e^{i\omega s} g_k(s) ds \right) d\omega$ for $j, k = 1, \dots, q+1$.

Proof. Follows from Theorem 1, Lemma 1 and the Cramer-Wold device via

$$\left| \int_0^H e^{-i\omega s} \left(\sum_{j=1}^{q+1} \lambda_j g_j(s) \right) ds \right|^2 = \sum_{j,k=1}^{q+1} \lambda_j \lambda_k \left(\int_0^H e^{-i\omega s} g_j(s) ds \right) \left(\int_0^H e^{i\omega s} g_k(s) ds \right).$$

■

8.2 Density of (X^s, Y^s) and Related Results

Let $Z = (X', Y)'$ and $U = \sqrt{X'X}$. Write μ_l for Lebesgue measure on $\mathbb{R}^l$, and ν_q for the surface measure of a q dimensional unit sphere. For $x \in \mathbb{R}^q$, let $x = x^s u$, where x^s is a point on the surface of a q dimensional unit sphere, and $u \in \mathbb{R}^+$. By Theorem 2.1.13 of Muirhead (1982), $d\mu_q(x) = u^{q-1} d\nu_q(x^s) d\mu_1(u)$. Further, for $y \in \mathbb{R}$, consider the change of variable $y = y^s u$ with $u \in \mathbb{R}^+$ and $y^s \in \mathbb{R}$, so that $d\mu_1(y) = u d\mu_1(y^s)$. We thus can write the joint density of (X^s, Y^s, U) with respect to $\nu_q \times \mu_1 \times \mu_1$ as

$$(2\pi)^{-(q+1)/2} |\Sigma|^{-1/2} \exp\left[-\frac{1}{2} \begin{pmatrix} x^s u \\ y^s u \end{pmatrix}' \Sigma^{-1} \begin{pmatrix} x^s u \\ y^s u \end{pmatrix}\right] u^q$$

and the marginal density of $Z^s = (X^{s'}, Y^s)'$ with respect to $\nu_q \times \mu_1$ is

$$\begin{aligned} f_{Z^s}(z^s) &= (2\pi)^{-(q+1)/2} |\Sigma|^{-1/2} \int_0^\infty u^q \exp\left[-\frac{1}{2} u^2 (z^{s'} \Sigma^{-1} z^s)\right] d\mu_1(u) \\ &= (2\pi)^{-(q+1)/2} |\Sigma|^{-1/2} \frac{1}{2} \int_0^\infty t^{(q-1)/2} \exp\left[-\frac{1}{2} t (z^{s'} \Sigma^{-1} z^s)\right] d\mu_1(t) \\ &= (2\pi)^{-(q+1)/2} |\Sigma|^{-1/2} \frac{1}{2} \Gamma\left(\frac{q+1}{2}\right) 2^{(q+1)/2} (z^{s'} \Sigma^{-1} z^s)^{-(q+1)/2} \\ &= \frac{1}{2} \pi^{-(q+1)/2} |\Sigma|^{-1/2} \Gamma\left(\frac{q+1}{2}\right) (z^{s'} \Sigma^{-1} z^s)^{-(q+1)/2} \end{aligned}$$

where the second equality follows from the form of the Gamma density function, and Γ denotes the gamma function. The implied marginal density of X^s is

$$f_{X^s}(x^s) = \frac{1}{2} \pi^{-(q)/2} |\Sigma_X|^{-1/2} \Gamma\left(\frac{q}{2}\right) (x^{s'} \Sigma_X^{-1} x^s)^{-q/2}.$$

Similarly, with $g(x^s) = E[\sqrt{X'X}|X^s = x^s]$, we obtain

$$\begin{aligned} f_{x^s}(x^s)g(x^s) &= \int_0^\infty u f_{(X^s, U)}(x^s, u) d\mu_1(u) \\ &= u(2\pi)^{-q/2} |\Sigma_{XX}|^{-1/2} \int_0^\infty u^{q-1} \exp[-\frac{1}{2}u^2(x^{s'}\Sigma_{XX}^{-1}x^s)] d\mu_1(u) \\ &= 2^{-1/2} \pi^{-q/2} |\Sigma_{XX}|^{-1/2} \Gamma(\frac{q+1}{2}) (x^{s'}\Sigma_{XX}^{-1}x^s)^{-(q+1)/2}. \end{aligned}$$

Finally, from (4), $\tilde{Y} = Y - \Sigma_{YX}\Sigma_{XX}^{-1}X \sim \mathcal{N}(0, \Sigma_{YY} - \Sigma_{YX}\Sigma_{XX}^{-1}\Sigma_{XY})$ and X are independent normal random variables. Also, using well known properties of a multivariate standard normal distribution, $X'\Sigma_{XX}^{-1}X \sim \chi_q^2$ is independent of $\tilde{X}^s = \Sigma_{XX}^{-1/2}X/\sqrt{X'\Sigma_{XX}^{-1}X}$. Since X^s is a one-to-one transformation of $\tilde{X}^s$, we thus obtain

$$\frac{\tilde{Y}}{\sqrt{X'\Sigma_{XX}^{-1}X/q} \sqrt{\Sigma_{YY} - \Sigma_{YX}\Sigma_{XX}^{-1}\Sigma_{XY}}} |X^s \sim t^q$$

and the result (6) follows by dividing the numerator and denominator by $\sqrt{X'X}$.

8.3 Approximate Least Favorable Distributions

In practice, it won't be possible to compute a least favorable distribution $\Lambda^\dagger$ that perfectly solves the program (12)-(14). To make further progress, we follow Elliott, Müller, and Watson (2014) (EMW in the following), and first formally state a lower bound on (12), and then define an approximate least favorable distribution (ALFD) Λ^* that solves (8) within a tolerance of ϵ .

To ease notation, write $V_W(A) = \int V_\theta(A) dW(\theta)$ and $C_\theta(A) = P_\theta(Y^s \in A(X^s))$. Also, we make the dependence of the set (16) on cv explicit by writing

$$A_{\Lambda, cv}(x^s) = \left\{ y^s : \frac{\int f_{(Y^s, X^s)|\theta}(y^s, x^s) d\Lambda(\theta)}{\int g_\theta(x^s) f_{X^s|\theta}(x^s) dW(\theta)} > cv \right\}. \quad (27)$$

We begin by proving the optimality of the set $A_{\Lambda, cv}$ in the problem $\min_A V_W(A)$ subject to $\int C_\theta(A) d\Lambda(\theta) = 1 - \alpha$.

Lemma 2 *Let $A_{\Lambda, cv}$ be such that $\int C_\theta(A_{\Lambda, cv}) d\Lambda(\theta) = 1 - \alpha$. Then $A_{\Lambda, cv}$ solves $\min_A V_W(A)$ subject to $\int C_\theta(A) d\Lambda(\theta) \geq 1 - \alpha$.*

Proof. Note that any A is equivalently characterized by the test-function $\varphi : \mathbb{R}^q \times \mathbb{R} \mapsto \{0, 1\}$ defined via $\varphi(y^s, x^s) = \mathbf{1}[y^s \in A(x^s)]$. In this notation, $V_W(A) = \int \int \int g_\theta(x^s) f_{X^s|\theta}(x^s) \varphi(y^s, x^s) d\nu_q(x^s) d\mu_1(y^s) dW(\theta) = \int \varphi(z^s) f_1(z^s) d\lambda_{q,1}(z^s)$, and $\int C_\theta(A) d\Lambda(\theta) = \int \int \int f_{Z^s|\theta}(x^s, y^s) \varphi(y^s, x^s) d\nu_q(x^s) d\mu_1(y^s) d\Lambda(\theta) = \int \varphi(z^s) f_0(z^s) d\lambda_{q,1}(z^s)$, where $d\lambda_{q,1}(z^s) = d\nu_q(x^s) \times d\mu_1(y^s)$, $f_1(z^s) = \int g_\theta(x^s) f_{X^s|\theta}(x^s) dW(\theta)$ and $f_0(z^s) = \int f_{Z^s|\theta}(z^s) d\Lambda(\theta)$. Thus, the

problem is equivalent to the problem of finding the best test that rejects (that is $\varphi = 1$) with probability at least $1 - \alpha$ when the “density” of Z^s is f_0 , and minimizes the rejection probability when the “density” of Z^s is f_1 . These densities do not necessarily integrate to one, but the solution still has to be of the Neyman-Pearson form (16), as can be seen by the very argument that proves the Neyman-Pearson Lemma: Set $\varphi^*(y^s, x^s) = \mathbf{1}[y^s \in A_{\Lambda, \text{cv}}(x^s)]$ and $\varphi(y^s, x^s) = \mathbf{1}[y^s \in A(x^s)]$ for some A that satisfies $\int C_\theta(A) d\Lambda(\theta) \geq 1 - \alpha$. Then $\int \varphi f_0 d\lambda_{q,1} \geq 1 - \alpha$ (we drop z^s as the dummy variable of integration for notational convenience), and

$$\begin{aligned} 0 &\leq \int (\varphi^* - \varphi)(f_0 - \text{cv } f_1) d\lambda_{q,1} \\ &\leq \text{cv} \left(\int \varphi f_1 d\lambda_{q,1} - \int \varphi^* f_1 d\lambda_{q,1} \right) \end{aligned}$$

where the first inequality follows from the definition of φ^* and the second from $1 - \alpha = \int \varphi^* f_0 d\lambda_{q,1} \leq \int \varphi f_0 d\lambda_{q,1}$. ■

A second result mirrors Lemma 1 of EMW and bounds the value of $\min_A V_W(A)$, formalizing the result verbally stated in Section 3.3.

Lemma 3 *Let $A_{\Lambda, \text{cv}}$ as in Lemma 2. Then for any A that satisfies $\inf_\theta C_\theta(A) \geq 1 - \alpha$, $V_W(A) \geq V_W(A_{\Lambda, \text{cv}})$.*

Proof. The result is immediate from Lemma 2 after noting that $\inf_\theta C_\theta(A) \geq 1 - \alpha$ implies $\int C_\theta(A) d\Lambda(\theta) \geq 1 - \alpha$. ■

Lemma 3 is useful, as it provides a set of lower bounds (indexed by Λ) on the achievable values of the objective (12). Thus, if a Λ can be identified that implies a small lower bound in the sense that a small adjustment to the critical value yields a set with uniform coverage and only marginally larger objective, the problem has been solved as a practical matter. Again following EMW, we denote such a distribution an ALFD.

Definition 2 *An ϵ -approximate least favorable distribution Λ^* is a probability distribution on θ satisfying*

- (i) *there exists cv^* such that $A_{\Lambda^*, \text{cv}^*}$ satisfies $\int C_\theta(A_{\Lambda^*, \text{cv}^*}) d\Lambda^*(\theta) = 1 - \alpha$*
- (ii) *there exists $\text{cv}^{*\epsilon} < \text{cv}^*$ such that $\inf_\theta \int C_\theta(A_{\Lambda^*, \text{cv}^{*\epsilon}}) \geq 1 - \alpha$, and $V_W(A_{\Lambda^*, \text{cv}^{*\epsilon}}) \leq V_W(A_{\Lambda^*, \text{cv}^*}) + \epsilon$.*

The strategy is thus to set some small tolerance level ϵ , and to numerically identify an ϵ -ALFD Λ^* . By definition, $A_{\Lambda^*, \text{cv}^{*\epsilon}}$ controls coverage uniformly, and invoking Lemma 3, its W -weighted average length is at most ϵ larger than of any prediction set that controls coverage uniformly.

Generalizations of Lemmas 2 and 3 for $A(x^s)$ additionally restricted to be a superset of some given set $B(x^s)$ are proven entirely analogously and are omitted for brevity (cf. Lemma 3 in EMW and Theorem 4 in Müller and Norets (2012)).

8.4 Computation of Frequentist Prediction Sets

8.4.1 Numerical Determination of $C_\theta(A)$ and $V_\theta(A)$

The algorithm for determining A_{MN} below requires repeated evaluation of $V_\theta(A)$ and $C_\theta(A)$, with A some set valued function of x^s . We employ an importance sampling Monte Carlo scheme: Write f_θ for the density of $Z^s = (X^s, Y^s)$ under θ , i.e. $f_\theta(z^s) = f_{(X^s, Y^s)|\theta}(x^s, y^s)$, and let f_p be a proposal density for Z^s such that f_θ is absolutely continuous with respect to f_p for all θ . Then

$$C_\theta(A) = E_p\left[\frac{f_\theta(Z^s)}{f_p(Z^s)}\mathbf{1}[Y^s \in A(X^s)]\right]$$

with E_p denoting integration with respect to f_p . Furthermore, from

$$\begin{aligned} V_\theta(A) &= E_\theta[g_\theta(X^s) \text{vol}(A(X^s))] \\ &= E_\theta\left[g_\theta(X^s) \int \mathbf{1}[y^s \in A(X^s)] dy^s\right] \\ &= \int \int g_\theta(x^s) f_{X^s|\theta}(x^s) \mathbf{1}[y^s \in A(x^s)] dx^s dy^s \end{aligned}$$

we obtain

$$V_\theta(A) = E_p\left[\frac{g_\theta(X^s) f_{X^s|\theta}(X^s)}{f_p(Z^s)} \mathbf{1}[Y^s \in A(X^s)]\right].$$

With N i.i.d. draws $Z_1^s, \dots, Z_m^s$ from f_p , we thus obtain the importance sampling approximations

$$\hat{C}_\theta(A) = N^{-1} \sum_{l=1}^N \frac{f_\theta(Z_l^s)}{f_p(Z_l^s)} \mathbf{1}[Y_l^s \in A(X_l^s)] \quad (28)$$

$$\hat{V}_\theta(A) = N^{-1} \sum_{l=1}^N \frac{g_\theta(X_l^s) f_{X^s|\theta}(X_l^s)}{f_p(Z_l^s)} \mathbf{1}[Y_l^s \in A(X_l^s)]. \quad (29)$$

The expression (29) is numerically advantageous, as it does not require any explicit numerical determination of the length of a set $A(x^s)$.

8.4.2 Numerical Determination of an ALFD

Our algorithm is analogous to what is suggested in EMW. It consists of two main parts: (I) determination of a candidate ϵ -ALFD Λ^* ; (II) numerical check of $\inf_{\theta \in \Theta} P_\theta(Y^s \in A_{\Lambda^*, \text{cv}^* \epsilon}(X^s)) \geq 1 - \alpha$. Both parts are based on a discrete grid of values for θ , with the grid Θ_{II} for part (II) finer and with a wider range than the grid Θ_I for part (I). In particular, with $\Delta_I = \{-0.4, -0.2, \dots, 1.0\}$, $\Delta_{II} = \{-0.4, -0.3, \dots, 1.0\}$, $\tilde{B}_I = \{0\} \cup \{e^{-5}, e^{-4}, \dots, e^5\}$, $\tilde{B}_{II} = \{0\} \cup \{e^{-5.5}, e^{-5.0}, \dots, e^{7.5}\}$, $C_I = \{0\} \cup \{e^{-3.0}, e^{-2.3}, \dots, e^{4.0}\}$, $C_{II} = \{0\} \cup \{e^{-3.35}, e^{-3.0}, \dots, e^{7.5}\}$, $\Theta_I = \{(b, 0, d)' : d \in \Delta_I, (8\pi)^{2db^2} \in \tilde{B}_I\} \cup \{(0, c, d)' : d \in \Delta_I, c \in C_I\} \cup \{(b, c, 1.0)' : c \in C_I, ((8\pi)^2 + c^2)^{db^2} \in \tilde{B}_I\}$ and $\Theta_{II} = \{(b, c, d)' : c \in C_{II}, d \in \Delta_{II}, ((8\pi)^2 + a^2)^{db^2} \in \tilde{B}_{II}\}$. Also let $\Theta_\Gamma = \{\theta : \theta = (0, 0, d), d \in \{-0.4, -0.2, \dots, 1.0\}\}$ denote the support of the prior Γ .

The range of values for b and c in grid Θ_{II} are chosen to approximately cover all possible values of Σ in the bcd -model (as $b \rightarrow \infty$ or $c \rightarrow \infty$, Σ converges to its value in the $I(0)$ model, for any value of $d \in \Delta_{II}$). Preliminary results suggested that the ALFD puts no mass on θ with $b > 0$, $c > 0$ and $d < 1$, motivating our smaller choice for Θ_I .

With Θ_I the support of candidate ALFDs Λ , and W putting weight w_θ on $\theta \in \Theta_I$, (27) becomes

$$A_{\Lambda, \text{cv}}(x^s) = \left\{ y^s : \frac{\sum_{\theta \in \Theta_I} \lambda_\theta f_{(Y^s, X^s)|\theta}(y^s, x^s)}{\sum_{\theta \in \Theta_I} w_\theta g_\theta(x^s) f_{X^s|\theta}(x^s)} > \text{cv} \right\}$$

where $\lambda_\theta \geq 0$ and $\sum_{\theta \in \Theta_I} \lambda_\theta = 1$. Note that $A_{\Lambda, \text{cv}}$ is fully determined by the ratios $\kappa_\theta = \lambda_\theta / \text{cv}$, $\theta \in \Theta_I$, so introduce the notation $A_{\Lambda, \text{cv}} = A_\kappa$.

Part (I) of the algorithm seeks to determine appropriate values of λ_θ , $\theta \in \Theta_I$, for an ϵ -ALFD. The following steps are for the A^{MN} set that also enforces the Bayes superset constraint (11).

1. Generate i.i.d. draws Z_l^s , $l = 1, \dots, N$ from $f_p(z^s) = |\Theta_I|^{-1} \sum_{\theta \in \Theta_I} f_\theta(z^s)$ (that is, the proposal density f_p is the equal probability mixture of f_θ with θ drawn uniformly from Θ_I).
2. Compute and store $f_\theta(Z_l^s)$ and $f_p(Z_l^s)$ for $\theta \in \Theta_I$, and $g_\theta(X_l^s) f_{X^s|\theta}(X_l^s)$ for $\theta \in \Theta_I$, $l = 1, \dots, N$.
3. Compute $V_\theta^{\text{known}} = \hat{V}_\theta(A_\theta^{\text{known}})$, $\theta \in \Theta_I$ where A_θ^{known} is the prediction set for known θ . To be specific, $A_\theta^{\text{known}} = [\mu_\theta(x^s) - t_{(1-\alpha/2)}^q \sigma_\theta(x^s), \mu_\theta(x^s) + t_{(1-\alpha/2)}^q \sigma_\theta(x^s)]$ with $\mu_\theta(x^s) = \Sigma_{YX} \Sigma_{XX}^{-1} x^s$ and $\sigma_\theta^2(x^s) = (\Sigma_{YY} - \Sigma_{YX} \Sigma_{XX}^{-1} \Sigma_{XY}) x^{s'} \Sigma_{XX}^{-1} x^s / q$, with Σ_{XX} , Σ_{YX} , Σ_{XY} and Σ_{YY} as implied by θ . Set $w_\theta = (V_\theta^{\text{known}})^{-1} / \sum_{t \in \Gamma_\theta} (V_t^{\text{known}})^{-1}$, $\theta \in \Gamma_\theta$.
4. For each Z_l^s , determine $\mathbf{1}[Y_l^s \in A^{\text{Bayes}}(X_l^s)] = \mathbf{1}[F^{\text{Bayes}}(Y_l^s | X_l^s) \in (\alpha/2, 1 - \alpha/2)]$, where F^{Bayes} is the c.d.f. of the Bayes predictive distribution, that is

$$F^{\text{Bayes}}(Y^s | X^s) = \frac{\sum_{\theta \in \Theta_I} \gamma_\theta f_{X^s|\theta}(X^s) F_t^q \left(\frac{Y^s - \mu_\theta(X^s)}{\sigma_\theta(X^s)} \right)}{\sum_{\theta \in \Theta_I} \gamma_\theta f_{X^s|\theta}(X^s)}$$

with F_t^q the c.d.f. of a student-t distribution with q degrees of freedom and $\gamma_\theta = 1/|\Theta_I|$.

5. Set $\kappa^{(0)} = (1, \dots, 1) \in \mathbb{R}^{|\Theta_I|}$, $\omega^{(0)} = (4, \dots, 4) \in \mathbb{R}^{|\Theta_I|}$, $\text{RP}^{(0)} = (\alpha, \dots, \alpha) \in \mathbb{R}^{|\Theta_I|}$
6. Iterate over $i = 1, \dots, 500$
 - (a) Compute $\text{RP}_\theta^{(i)} = 1 - \hat{C}_\theta(A_{\kappa^{(i)}} \cup A^{\text{Bayes}})$, $\theta \in \Theta_I$.
 - (b) Set $\kappa_\theta^{(i+1)} = \kappa_\theta^{(i)} \exp(\omega_\theta^{(i)} (\text{RP}_\theta^{(i)} - \alpha))$, $\theta \in \Theta_I$.
 - (c) If $i \geq 400$, set $\omega_\theta^{(i+1)} = 0.1$. Otherwise, set $\omega_\theta^{(i+1)} = 1.03 \omega_\theta^{(i)}$ if $(\text{RP}_\theta^{(i)} - \alpha)(\text{RP}_\theta^{(i)} - \alpha) > 0$ and $\omega_\theta^{(i+1)} = 0.5 \omega_\theta^{(i)}$ otherwise, $\theta \in \Theta_I$.
7. Set $\lambda_\theta^* = \kappa_\theta^{(500)} / \sum_{t \in \Theta_I} \kappa_t^{(500)}$, $\theta \in \Theta_I$.

As in EMW, the simple idea underlying Step 6.b is a fixed-point iteration that decreases κ_θ under overcoverage, and increases it otherwise, with the change (in logarithms) proportional to the degree of over- or undercoverage. The idea of Step 6.c is to half the step length if the sign of the coverage violation switched in the last iteration, and to gradually increase it otherwise. In the last 100 of the 500 iterations, a small step length is enforced regardless.

One would expect λ_θ^* of Step 7 to be a rough approximation to the least favorable distribution with the parameter space restricted to Θ_I . With Θ_I a sufficiently fine grid, this in turn is a candidate for an ϵ -ALFD Λ^* in the sense of Definition 2. Part (II) of the algorithm checks whether Λ^* constructed from λ_θ^* does indeed satisfy the properties of an ϵ -ALFD Λ^* .

1. Compute the number cv^* such that $\sum_{\theta \in \Theta_I} w_\theta \hat{C}_\theta(A_{\Lambda^*, \text{cv}^*} \cup A^{\text{Bayes}}) = 1 - \alpha$.
2. Compute the number cv^{ϵ} such that $\sum_{\theta \in \Theta_I} w_\theta \hat{V}_\theta(A_{\Lambda^*, \text{cv}^{\epsilon}} \cup A^{\text{Bayes}}) = \epsilon + \sum_{\theta \in \Theta_I} w_\theta \hat{V}_\theta(A_{\Lambda^*, \text{cv}^*} \cup A^{\text{Bayes}})$.
3. Check whether $\inf_{\theta \in \Theta_{II}} \hat{C}_\theta(A_{\Lambda^*, \text{cv}^{\epsilon}} \cup A^{\text{Bayes}}) \geq 1 - \alpha$.

We set $\epsilon = 0.01$ throughout, so that in terms of the W -average expected length, the set $A^{MN} = A_{\Lambda^*, \text{cv}^{\epsilon}} \cup A^{\text{Bayes}}$ is at most 1% longer than any set satisfying (9)-(11). We use $N = 250,000$, leading to Monte Carlo errors of $\hat{C}_\theta$ of approximately 0.1-0.3%. For a given value of q , α and r , these computations take approximately one minute on a modern PC using Fortran. We compute A^{MN} sets for $q \in \{6, 12, 24, 48\}$ and $r \in \{0.075, 0.1, 0.15\} \cup \{0.2, 0.3, \dots, 1.0\} \cup \{1.2, 1.4, 1.6, 1.8\}$. In the empirical work, prediction intervals for other values of r are computed via linear interpolation.

8.4.3 Computation of Σ in bcd -model

One possibility to compute Σ is to notice that for all g_j under consideration, $\int_0^{1+r} g_j(s) e^{i\omega s} ds$ can be obtained in closed form, so that one can apply one-dimensional numerical integration to (7) for $S(\omega) = (c^2 + \omega^2)^{-d} + b^2$. It turns out, however, that the integrand is highly oscillatory and, for d small ($d = -0.4$), very slowly decaying.

We therefore now derive a numerically more stable expression for $\Sigma_{j,k}$ in the bcd -model. Note that (7) is linear in S , and Σ for constant S can be computed from (7) in closed form. It thus suffices so obtain a more convenient expression for $b = 0$ (since the additional component stemming from the $I(0)$ component for $b > 0$ can be added easily).

Now for $c > 0$ and $1/2 < d$, from (4.58) in Lindgren (2013)

$$\gamma_{c,d}(s) = \int_{-\infty}^{\infty} (c^2 + \omega^2)^{-d} e^{i\omega s} d\omega = \frac{2^{\frac{3}{2}-d} \sqrt{\pi} (|s|/c)^{d-1/2} I_{d-1/2}(c|s|)}{\Gamma(d)} \quad (30)$$

where $I_\nu(x)$ is the modified Bessel function of the second kind, so that

$$\Sigma_{j,k} = \int_{-\infty}^{\infty} \int_0^{1+r} \int_0^{1+r} g_j(s) e^{i\omega s} g_k(u) e^{-i\omega u} (c^2 + \omega^2)^{-d} e^{i\omega s} d\omega ds du \quad (31)$$

$$= \int_0^{1+r} \int_0^{1+r} g_j(s) g_k(u) \gamma_{c,d}(s-u) ds du. \quad (32)$$

We approximate (32) by a double sum with a step length of 0.0002. For $c = 0$, we use the implied value for $c = 0.01$.

It remains to obtain a suitable expression for $-1/2 < d < 1/2$. Suppose g is differentiable on (H_L, H_U) with derivative g' . Then by integration by parts,

$$\begin{aligned} \int_{H_L}^{H_U} g(s) e^{i\omega s} ds &= \int_{H_L}^{H_U} g(s) e^{-cs} e^{(c+i\omega)s} ds \\ &= g(H_U) e^{-H_U c} \frac{e^{(c+i\omega)H_U}}{c+i\omega} - g(H_L) \frac{1}{c+i\omega} - \int_{H_L}^{H_U} (g'(s) e^{-cs} - c g(s) e^{-cs}) \frac{e^{(c+i\omega)s}}{c+i\omega} ds \\ &= \frac{g(H_U) e^{i\omega H_U} - g(H_L) e^{i\omega H_L} - \int_{H_L}^{H_U} (g'(s) - c g(s)) e^{i\omega s} ds}{c+i\omega}. \end{aligned}$$

All g functions entering $\Sigma_{j,k}$ are differentiable on $(0, 1)$ and on $(1, 1+r)$, so we obtain

$$\int_0^{1+r} g(s) e^{i\omega s} ds = \frac{g(1+r) e^{i\omega(1+r)} - (g^+(1) - g(1)) e^{i\omega} - g(0) - \int_0^{1+r} (g'(s) - c g(s)) e^{i\omega s} ds}{c+i\omega} \quad (33)$$

where $g^+(1)$ the right limit of $g(s)$ at $s = 1$, and $g'(s)$ for $s \in \{0, 1, 1+r\}$ may be defined arbitrarily. Now applying (33) in (31) twice, and noting that $(c+i\omega)(c-i\omega) = c^2 + \omega^2$, all integrals with respect to $d\omega$ in (31) are of the form $C_1 \int_{-\infty}^{\infty} (c^2 + \omega^2)^{-(d+1)} e^{i\omega(s-C_2)} d\omega$ for C_1 and C_2 not depending on ω , which can be computed in closed form using (30). We obtain

$$\begin{aligned} \Sigma_{j,k} &= \begin{pmatrix} g_j(1+r) \\ g_j(1) - g_j^+(1) \\ -g_j(0) \end{pmatrix}' \begin{pmatrix} \gamma_{c,d+1}(0) & \gamma_{c,d+1}(r) & \gamma_{c,d+1}(1+r) \\ \gamma_{c,d+1}(r) & \gamma_{c,d+1}(0) & \gamma_{c,d+1}(1) \\ \gamma_{c,d+1}(1+r) & \gamma_{c,d+1}(1) & \gamma_{c,d+1}(0) \end{pmatrix} \begin{pmatrix} g_k(1+r) \\ g_k(1) - g_k^+(1) \\ -g_k(0) \end{pmatrix} \\ &\quad - \begin{pmatrix} g_j(1+r) \\ g_j(1) - g_j^+(1) \\ -g_j(0) \end{pmatrix}' \begin{pmatrix} \int_0^{1+r} (g'_k(s) - c g_k(s)) \gamma_{c,d+1}(1+r-s) ds \\ \int_0^{1+r} (g'_k(s) - c g_k(s)) \gamma_{c,d+1}(1-s) ds \\ \int_0^{1+r} (g'_k(s) - c g_k(s)) \gamma_{c,d+1}(s) ds \end{pmatrix} \\ &\quad - \begin{pmatrix} g_k(1+r) \\ g_k(1) - g_k^+(1) \\ -g_k(0) \end{pmatrix}' \begin{pmatrix} \int_0^{1+r} (g'_j(s) - c g_j(s)) \gamma_{c,d+1}(1+r-s) ds \\ \int_0^{1+r} (g'_j(s) - c g_j(s)) \gamma_{c,d+1}(1-s) ds \\ \int_0^{1+r} (g'_j(s) - c g_j(s)) \gamma_{c,d+1}(s) ds \end{pmatrix} \\ &\quad + \int_0^{1+r} \int_0^{1+r} (g'_j(s) - c g_j(s)) (g'_k(u) - c g_k(u)) \gamma_{c,d+1}(s-u) ds du. \end{aligned}$$

We again approximate the integrals in this expression by sums with a step length of 0.0002, and approximate Σ for $c = 0$ by the value implied by $c = 0.01$.

8.5 Proof of Theorem 1

We now introduce some notation that is used in the proof of Theorem 1, and in auxiliary Lemmas. We first state these Lemmas and then the proof of Theorem 1.

Notation: Define $\sigma_T^2 = \text{Var}[T^{-1/2} \int_0^H g(s)x_{T, \lfloor sT \rfloor + 1} ds]$. Note that with $\tilde{g}_{T,t} = T \int_{(t-1)/T}^{t/T} g(s) ds$ and $\tilde{G}_{T,t} = T^{-1} \sum_{s=1}^{t-1} \tilde{g}_{T,t} = \sum_{s=1}^{t-1} \int_{(s-1)/T}^{s/T} g(s) ds = \int_0^{(t-1)/T} g(s) ds = G(\frac{t-1}{T})$, we find using summation by parts

$$\begin{aligned} T^{-1/2} \int_0^H g(s)x_{T, \lfloor sT \rfloor + 1} ds &= T^{-3/2} \sum_{t=1}^{HT} \tilde{g}_{T,t} x_{T,t} \\ &= T^{-1/2} \tilde{G}_{T, HT+1} x_{T, HT} - T^{-1/2} \sum_{t=1}^{HT} \tilde{G}_{T,t} \Delta x_{T,t} \\ &= -T^{1/2} \sum_{t=1}^{HT} G(\frac{t-1}{T}) \Delta x_{T,t} \end{aligned}$$

since $\tilde{G}_{T, HT+1} = G(H) = 0$.

Lemma 4 $\sigma_T^2 \rightarrow \sigma^2 = \int_{-\infty}^{\infty} S(\omega) \left| \int_0^H e^{-i\omega s} g(s) ds \right|^2 d\omega$.

Proof. Let γ_T be the autocovariances of $\Delta x_{T,t}$. We find

$$\begin{aligned} \text{Var}[T^{-1/2} \sum_{t=1}^{HT} G(\frac{t-1}{T}) \Delta x_{T,t}] &= T^{-1} \sum_{j,k=1}^{HT} \gamma_T(k-j) G(\frac{k-1}{T}) G(\frac{j-1}{T}) \\ &= T^{-1} \sum_{j,k=1}^{HT} \left(\int_{-\pi}^{\pi} e^{i\lambda(k-j)} F_T(\lambda) d\lambda \right) G(\frac{k-1}{T}) G(\frac{j-1}{T}) \\ &= T^{-1} \int_{-\pi}^{\pi} F_T(\lambda) \left| \sum_{t=1}^{HT} e^{i\lambda t} G(\frac{t-1}{T}) \right|^2 d\lambda. \end{aligned}$$

Now for any fixed K , we will show

$$\begin{aligned} T^{-1} \int_{-K/T}^{K/T} F_T(\lambda) \left| \sum_{t=1}^{HT} e^{i\lambda t} G(\frac{t-1}{T}) \right|^2 d\lambda &= \int_{-K}^K F_T(\frac{\omega}{T}) \left| T^{-1} \sum_{t=1}^{HT} e^{i\omega t/T} G(\frac{t-1}{T}) \right|^2 d\omega \\ &\rightarrow \int_{-K}^K S(\omega) \omega^2 \left| \int_0^H e^{i\omega s} G(s) ds \right|^2 d\omega. \end{aligned} \quad (34)$$

First note that for any two complex numbers a, b , $|a - b| \geq ||a| - |b||$, so that $||a|^2 - |b|^2| = |(|a| + |b|)(|a| - |b|)| \leq (|a| + |b|)|a - b|$. Thus,

$$\begin{aligned} \left| \left| \int_0^H e^{i\omega s} G(s) ds \right|^2 - \left| T^{-1} \sum_{t=1}^{HT} e^{i\omega t/T} G(\frac{t-1}{T}) \right|^2 \right| &\leq \\ &2 \sup_s |G(s)| \cdot \left| T^{-1} \sum_{t=1}^{HT} e^{i\omega t/T} G(\frac{t-1}{T}) - \int_0^H e^{i\omega s} G(s) ds \right| \end{aligned}$$

and

$$\begin{aligned}
& \left| T^{-1} \sum_{t=1}^{HT} e^{i\omega t/T} G\left(\frac{t-1}{T}\right) - \int_0^H e^{i\omega s} G(s) ds \right| \\
& \leq \int_0^H |e^{i\omega s} G(s) - e^{i\omega [sT]/T} G([sT]/T)| ds \\
& = \int_0^H |e^{i\omega s} (G(s) - G([sT]/T)) + G([sT]/T) (e^{i\omega s} - e^{i\omega [sT]/T})| ds \\
& \leq \int_0^H |G(s) - G([sT]/T)| ds + \sup_s |G(s)| \int_0^H |1 - e^{i\omega([sT]/T - s)}| ds.
\end{aligned}$$

Since for any real a , $|1 - e^{ia}| < |a|$, $\sup_{|\omega| \leq K} |1 - e^{i\omega([sT]/T - s)}| \leq K/T$, and also $|G(s) - G([sT]/T)| \leq T^{-1} \sup_s |g(s)|$. Thus,

$$\sup_{|\omega| \leq K} \left| \left| T^{-1} \sum_{t=1}^{HT} e^{i\omega t/T} G\left(\frac{t-1}{T}\right) \right|^2 - \left| \int_0^H e^{i\omega s} G(s) ds \right|^2 \right| \rightarrow 0$$

and (34) follows from assumption (iii.a) by straightforward arguments.

Further, since $\int_0^H e^{i\omega s} g(s) ds = -i\omega \int_0^H e^{i\omega s} G(s) ds$,

$$\int_{-K}^K S(\omega) \omega^2 \left| \int_0^H e^{i\omega s} G(s) ds \right|^2 d\omega = \int_{-K}^K S(\omega) \left| \int_0^H e^{i\omega s} g(s) ds \right|^2 d\omega.$$

Thus, for any fixed K ,

$$\delta_K(T) = T^{-1} \int_{-K/T}^{K/T} F_T(\lambda) \left| \sum_{t=1}^{HT} e^{i\lambda t} G\left(\frac{t-1}{T}\right) \right|^2 d\lambda - \int_{-K}^K S(\omega) \left| \int_0^H e^{i\omega s} g(s) ds \right|^2 d\omega \rightarrow 0.$$

Now for each T , define K_T as the largest integer $K \leq T$ for which $\sup_{T' \geq T} |\delta_K(T')| \leq 1/K$ (and zero if no such K exists). Note that $\delta_K(T) \rightarrow 0$ for all fixed K implies that $K_T \rightarrow \infty$, and by construction, also $\delta_{K_T}(T) \rightarrow 0$. The result now follows from

$$T^{-1} \int_{K_T/T}^{\pi} F_T(\lambda) \left| \sum_{t=1}^{HT} e^{i\lambda t} G\left(\frac{t-1}{T}\right) \right|^2 d\lambda \leq C^2 T^{-3} \int_{K_T/T}^{\pi} F_T(\lambda) \frac{1}{\lambda^4} d\lambda \rightarrow 0$$

by assumptions (iii.b) and (iv). ■

Lemma 5 $T^{-1/2} \sup_{t,T} \left| \sum_{j=1}^{HT} G\left(\frac{j-1}{T}\right) c_{T,j-t} \right| \rightarrow 0$.

Proof. Recall that for any two real, square integrable sequences $\{a_j\}_{j=-\infty}^{\infty}$ and $\{b_j\}_{j=-\infty}^{\infty}$, $\sum_{j=-\infty}^{\infty} a_j b_j = \frac{1}{2\pi} \int_{-\pi}^{\pi} \hat{A}(\lambda) \hat{B}^*(\lambda) d\lambda$, where $\hat{A}(\lambda) = \sum_{j=-\infty}^{\infty} a_j e^{-i\lambda j}$ and $\hat{B}^*(\lambda) = \sum_{j=-\infty}^{\infty} b_j e^{i\lambda j}$. Thus

$$\sum_{j=1}^{HT} G\left(\frac{j-1}{T}\right) c_{T,j-t} = \sum_{j=1-t}^{HT-t} G\left(\frac{t+j-1}{T}\right) c_{T,j} = \frac{1}{2\pi} \int_{-\pi}^{\pi} e^{i\lambda t} \hat{G}_T(\lambda) \hat{C}_T^*(\lambda) d\lambda$$

where $\hat{C}_T(\lambda) = \sum_{j=-\infty}^{\infty} c_{T,j} e^{-i\lambda j}$ and $\hat{G}_T(\lambda) = \sum_{j=1}^{HT} e^{-i\lambda j} G(\frac{j-1}{T})$, so that $e^{i\lambda t} \hat{G}_T(\lambda) = \sum_{j=1}^{HT} e^{-i\lambda(j-t)} G(\frac{j-1}{T}) = \sum_{j=1-t}^{HT-t} e^{-i\lambda j} G(\frac{j+1-t}{T})$, and $\hat{C}_T^*$ is the complex conjugate of $\hat{C}_T$. Now since $F_T(\lambda) = \frac{1}{2\pi} |\hat{C}_T(\lambda)|^2$, we find by the Cauchy-Schwarz inequality that

$$2\pi \left| \sum_{j=1}^{HT} G(\frac{j-1}{T}) c_{T,j-t} \right| = \left| \int_{-\pi}^{\pi} e^{i\lambda t} \hat{G}_T(\lambda) \hat{C}_T^*(\lambda) d\lambda \right| \leq \sqrt{2\pi} \int_{-\pi}^{\pi} |\hat{G}_T(\lambda)| F_T(\lambda)^{1/2} d\lambda.$$

Also, since $|\hat{G}_T(\lambda)| \leq \sum_{j=1}^{HT} |G(\frac{j-1}{T})|$, we have

$$\begin{aligned} \int_0^{1/T} |\hat{G}_T(\lambda)| F_T(\lambda)^{1/2} d\lambda &\leq T^{-1} \sum_{j=1}^{HT} |G(\frac{j-1}{T})| \cdot \int_0^1 F_T(\omega/T)^{1/2} d\omega \\ &\rightarrow \int_0^H |G(s)| ds \cdot \int_0^1 \omega S(\omega)^{1/2} d\omega < \infty \end{aligned}$$

where the convergence follows from assumption (iii.a). Furthermore, by assumption (iv),

$$\int_{1/T}^{\pi} |\hat{G}_T(\lambda)| F_T(\lambda)^{1/2} d\lambda \leq CT^{-1} \int_{1/T}^{\pi} F_T(\lambda)^{1/2} \lambda^{-2} d\lambda.$$

The result follows by assumption (iii.c). ■

Lemma 6 *For every $\epsilon > 0$ there exists a $M > 0$ such that*

$$\text{Var}[T^{-1/2} \int_0^H g(s) x_{T, \lfloor sT \rfloor + 1} ds + T^{-1/2} \sum_{t=-MT}^{MT} \left(\sum_{j=1}^{HT} G(\frac{j-1}{T}) c_{T,j-t} \right) \varepsilon_t] < \epsilon.$$

For this M , $\sigma_{T,M}^2 = \text{Var}[T^{-1/2} \sum_{t=-MT}^{MT} \left(\sum_{j=1}^{HT} G(\frac{j-1}{T}) c_{T,j-t} \right) \varepsilon_t]$ satisfies $\limsup_{T \rightarrow \infty} |\sigma_{M,T}^2 - \sigma^2| < \epsilon$.

Proof. We have

$$\begin{aligned} - \int_0^H g(s) x_{T, \lfloor sT \rfloor + 1} ds &= \sum_{t=1}^{HT} G(\frac{t-1}{T}) \Delta x_{T,t} \\ &= \sum_{t=1}^{HT} G(\frac{t-1}{T}) \sum_{s=-\infty}^{\infty} c_{T,s} \varepsilon_{t-s} \\ &= \sum_{t=-\infty}^{\infty} \left(\sum_{j=1}^{HT} G(\frac{j-1}{T}) c_{T,j-t} \right) \varepsilon_t \end{aligned}$$

so that, with $\bar{G} = \sup_{0 \leq s < H} |G(s)|$

$$\text{Var}[T^{-1/2} \int_0^H g(s) x_{T, \lfloor sT \rfloor + 1} ds + T^{-1/2} \sum_{t=-MT}^{MT} \left(\sum_{j=1}^{HT} G(\frac{j-1}{T}) c_{T,j-t} \right) \varepsilon_t]$$

$$\begin{aligned}
&= T^{-1} \sum_{l=-\infty}^{-MT-1} \left(\sum_{j=1}^{HT} G\left(\frac{j-1}{T}\right) c_{T,j-l} \right)^2 + T^{-1} \sum_{l=MT+1}^{\infty} \left(\sum_{j=1}^{HT} G\left(\frac{j-1}{T}\right) c_{T,j-l} \right)^2 \\
&\leq \bar{G}^2 T^{-1} \sum_{l=MT+1}^{\infty} \left(\sum_{j=1}^{HT} (|c_{T,j-l}| + |c_{T,j+l}|) \right)^2 \\
&\leq 4H^2 \bar{G}^2 T^{-1} \sum_{l=MT+1}^{\infty} \left(T \sup_{|s| \geq l-HT} |c_{T,s}| \right)^2 \\
&= 4H^2 \bar{G}^2 T^{-1} \sum_{l=(M-H)T+1}^{\infty} \left(T \sup_{|s| \geq l} |c_{T,s}| \right)^2
\end{aligned}$$

which can be made arbitrarily small by choosing M large enough via assumption (ii).

The second claim follows directly from Lemma 4. ■

Lemma 7 For any large enough integer $M > 0$, $\sigma_{M,T}^{-1} T^{-1/2} \sum_{t=-MT}^{MT} \left(\sum_{j=1}^{HT} G\left(\frac{j-1}{T}\right) c_{T,j-t} \right) \varepsilon_t \Rightarrow \mathcal{N}(0, 1)$.

Proof. By the second claim in Lemma 6 and Lemma 4, $\sigma_{M,T} = O(1)$ and $\sigma_{M,T}^{-1} = O(1)$. By Theorem 24.3 in Davidson (1994), it thus suffices to show (a) $T^{-1/2} \sup_{1 \leq t \leq HT} \left| \sum_{j=1}^{HT} G\left(\frac{j-1}{T}\right) c_{T,j-t} \varepsilon_t \right| \xrightarrow{P} 0$ and (b) $T^{-1} \sum_{t=-MT}^{MT} \left(\sum_{j=1}^{HT} G\left(\frac{j-1}{T}\right) c_{T,j-t} \right)^2 (\varepsilon_t^2 - 1) \xrightarrow{P} 0$.

(a) is implied by the Lyapunov condition via Davidson's (1994) Theorems 23.16 and 23.11. Thus, it suffices to show that

$$\sum_{t=-MT}^{MT} E \left[\left| T^{-1/2} \left(\sum_{j=1}^{HT} G\left(\frac{j-1}{T}\right) c_{T,j-t} \right) \varepsilon_t \right|^{2+\delta} \right] \rightarrow 0.$$

Now

$$\begin{aligned}
&\sum_{t=-MT}^{MT} E \left[\left| T^{-1/2} \left(\sum_{j=1}^{HT} G\left(\frac{j-1}{T}\right) c_{T,j-t} \right) \varepsilon_t \right|^{2+\delta} \right] \\
&\leq T^{-1-\delta/2} \sum_{t=-MT}^{MT} \left| \sum_{j=1}^{HT} G\left(\frac{j-1}{T}\right) c_{T,j-t} \right|^{2+\delta} E(|\varepsilon_t|^{2+\delta}) \\
&\leq \left(\sup_t E(|\varepsilon_t|^{2+\delta}) \right) \cdot T^{-\delta/2} \sup_t \left| \sum_{j=1}^{HT} G\left(\frac{j-1}{T}\right) c_{T,j-t} \right|^\delta \cdot T^{-1} \sum_{t=-MT}^{MT} \left| \sum_{j=1}^{HT} G\left(\frac{j-1}{T}\right) c_{T,j-t} \right|^2 \\
&= \left(\sup_t E(|\varepsilon_t|^{2+\delta}) \right) \cdot \left(T^{-1/2} \sup_t \left| \sum_{j=1}^{HT} G\left(\frac{j-1}{T}\right) c_{T,j-t} \right| \right)^\delta \cdot \sigma_{M,T}^2 \rightarrow 0
\end{aligned}$$

where the convergence follows from Lemma 5.

For (b), we apply Theorem 19.11 of Davidson (1994) with Davidson's X and c chosen as $X_{T,t}^D = c_{T,t}^D(\varepsilon_t^2 - 1)$ and $c_{T,t}^D = T^{-1} \left(\sum_{j=1}^{HT} G(\frac{j-1}{T}) c_{T,j-t} \right)^2$. Then $X_{T,t}^D / c_{T,t}^D = \varepsilon_t^2 - 1$ is uniformly integrable, since $\sup_t E(|\varepsilon_t|^{2+\delta}) < \infty$. Further,

$$\sum_{t=-MT}^{MT} c_{T,t}^D = T^{-1} \sum_{t=-MT}^{MT} \left(\sum_{j=1}^{HT} G(\frac{j-1}{T}) c_{T,j-t} \right)^2 = \sigma_{M,T}^2 = O(1)$$

and

$$\begin{aligned} \sum_{t=-MT}^{MT} (c_{T,t}^D)^2 &= T^{-2} \sum_{t=-MT}^{MT} \left(\sum_{j=1}^{HT} G(\frac{j-1}{T}) c_{T,j-t} \right)^4 \\ &= \sup_{1 \leq t \leq HT} \left| T^{-1/2} \sum_{j=1}^{HT} G(\frac{j-1}{T}) c_{T,j-t} \right|^2 \cdot T^{-1} \sum_{t=-MT}^{MT} \left(\sum_{j=1}^{HT} G(\frac{j-1}{T}) c_{T,j-t} \right)^2 \\ &= \sup_{1 \leq t \leq HT} \left| T^{-1/2} \sum_{j=1}^{HT} G(\frac{j-1}{T}) c_{T,j-t} \right|^2 \cdot \sigma_{M,T}^2 \rightarrow 0 \end{aligned}$$

where the convergence follows from Lemma 5. ■

Proof of Theorem 1:

By Lemmas 4, 6 and 7, for large enough M ,

$$\sigma^{-1} T^{-1/2} \int_0^H g(s) x_{T, \lfloor sT \rfloor + 1} ds = \frac{\sigma_{M,T}}{\sigma} A_T + \frac{\sigma_{M,T}}{\sigma} B_T$$

where $A_T \Rightarrow \mathcal{N}(0, 1)$, $\limsup_{T \rightarrow \infty} E(B_T^2) < \epsilon$ and $\limsup_{T \rightarrow \infty} |\sigma_{M,T}^2 - \sigma^2| < \epsilon$. Thus, by Slutsky's Theorem, as $\epsilon \rightarrow 0$, the desired convergence in distribution follows. But ϵ was arbitrary, which proves the Theorem.

8.6 Autocovariances for the bcd -model

As described in the text, equation (18) is the local-to-zero spectrum of the process $x_t = e_{1t} + (bT^d)^{-1} z_t$ where $(1 - \rho_T L)^d z_t = e_{2t}$ with $\rho_T = 1 - c/T$ and e_{1t} and e_{2t} mutually uncorrelated white noise processes with common variance σ^2 . The autocovariances for the process are the sum of the autocovariances for its two components. The autocovariances for the second component can be computed as follows. Suppressing the dependence of ρ on T , the spectrum for z is

$$R(\lambda) = \frac{\sigma^2}{2\pi} (1 + \rho^2 - 2\rho \cos(\lambda))^{-d}$$

and the k 'th autocovariance is therefore

$$\gamma_k = \int_0^\pi \cos(k\lambda) R(\lambda) d\lambda$$

$$\begin{aligned}
&= \frac{\sigma^2}{2\pi} \int_0^\pi \cos(k\lambda) (1 + \rho^2 - 2\rho \cos(\lambda))^{-d} d\lambda \\
&= \frac{\sigma^2}{2\pi} \left(\frac{\rho}{(1 + \rho)^2} \right)^k (1 + \rho)^{-2d} \left(\frac{\Gamma(d + k)}{\Gamma(d)\Gamma(1 + k)} \right) {}_2F_1 \left(0.5 + k, d + k; 1 + 2k; \frac{4\rho}{(1 + \rho)^2} \right)
\end{aligned}$$

where ${}_2F_1(a, b; c, z)$ is the hypergeometric function.

From Luke (1975, page 271 equation (7)):

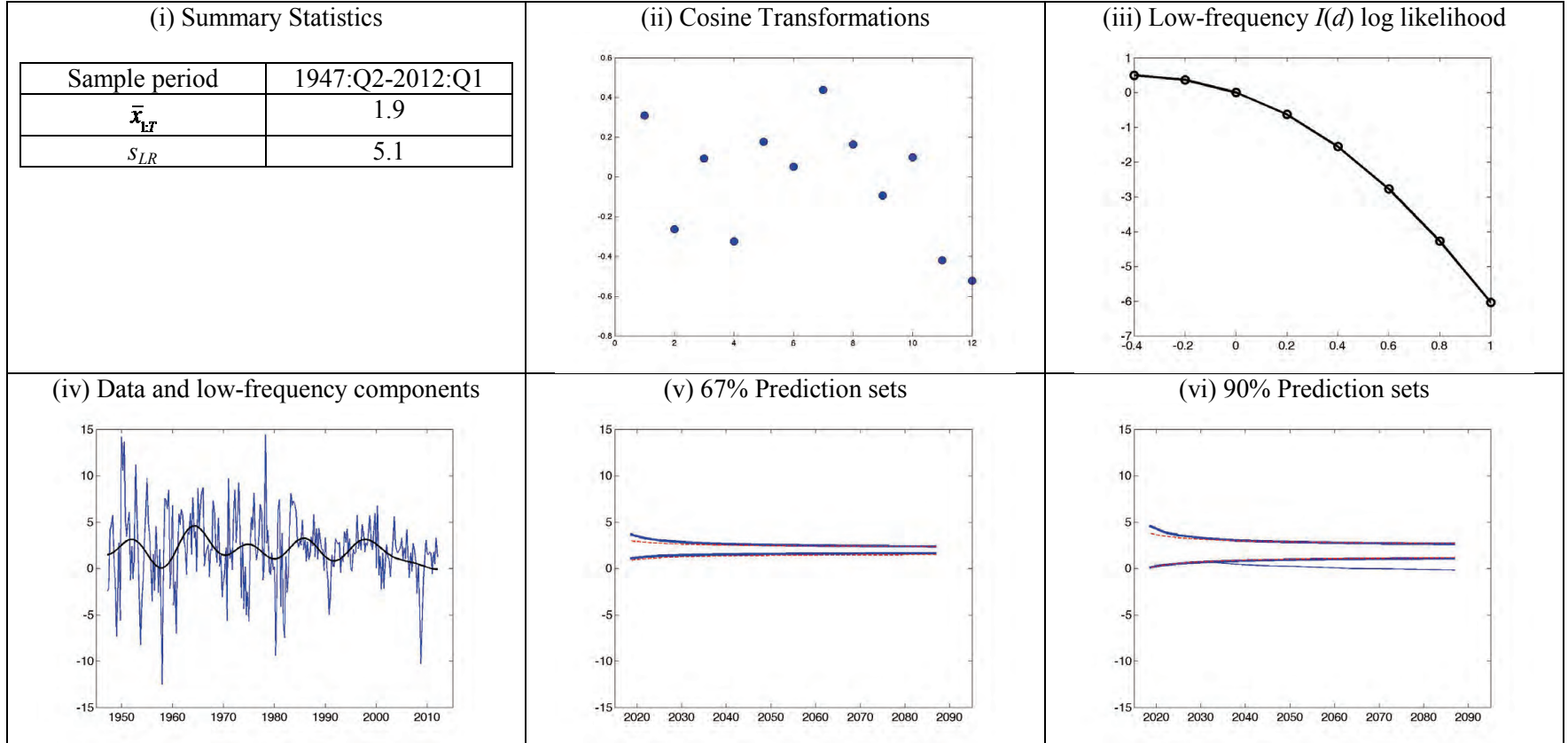
$${}_2F_1(a, b; 2a; z) = \left(\frac{2}{1 + y} \right)^{2b} {}_2F_1 \left(b, b + 0.5 - a; a + 0.5; \left(\frac{1 - y}{1 + y} \right)^2 \right), \text{ where } y = (1 - z)^{1/2}.$$

Thus, the expression for γ_k simplifies:

$$\begin{aligned}
\gamma_k &= \frac{\sigma^2}{2\pi} \left(\frac{\rho}{(1 + \rho)^2} \right)^k (1 + \rho)^{2(d+k)} (1 + \rho)^{-2d} \left(\frac{\Gamma(d + k)}{\Gamma(d)\Gamma(1 + k)} \right) {}_2F_1(d + k, d; 1 + k; \rho) \\
&= \frac{\sigma^2}{2\pi} \rho^k \left(\frac{\Gamma(d + k)}{\Gamma(d)\Gamma(1 + k)} \right) {}_2F_1(d + k, d; 1 + k; \rho^2).
\end{aligned}$$

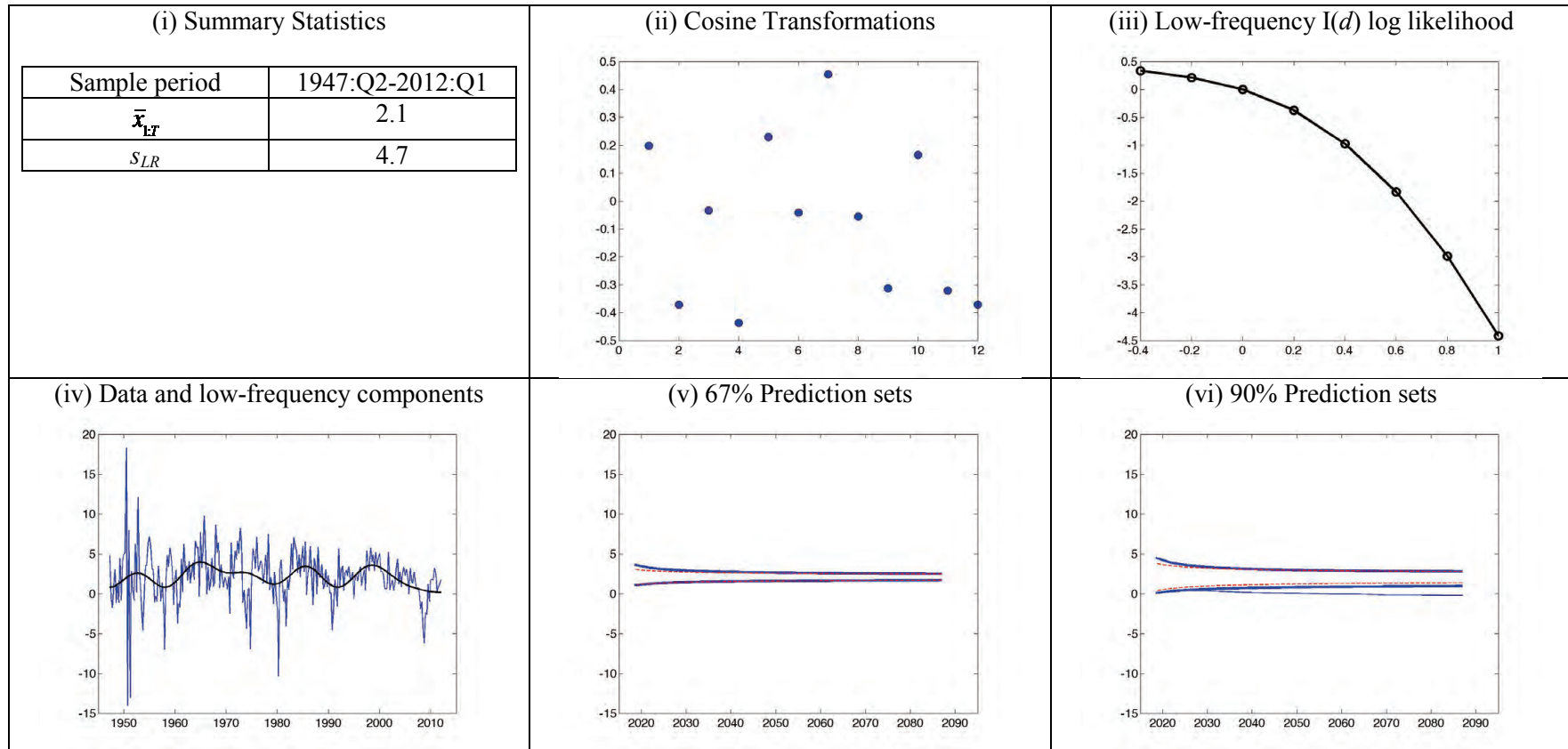
Additional Figures

A.1 Growth rate of real per-capita GDP



Notes: The cosine transformations shown in (ii) are the standardized values, $X_{T,1q}^*$. The low-frequency $I(d)$ likelihood is computed using $X_{T,1q}^*$ and its asymptotic distribution given in the text; values are relative to the $I(0)$ model. The low-frequency components in (iv) are the projection of the series onto $\cos[(t-0.5)\pi/T]$ for $j = 0, \dots, 12$. The predictions sets in panels (v) and (vi) are A^{Bayes} (thick line), A^{MN} (thin blue), and $A^{I(0)}$ (dashed red), and are computed separately for each horizon, so coverage levels hold pointwise, not jointly across horizons.

A.2 Growth rate of real per-capita consumption expenditures



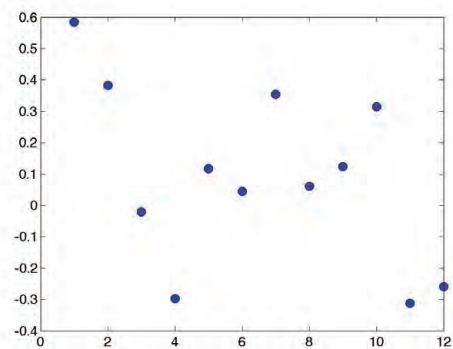
See notes to A.1.

A.3 Growth rate of total factor productivity

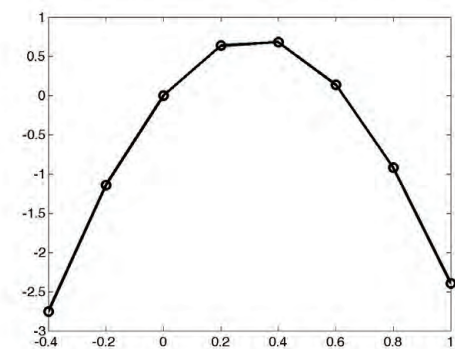
(i) Summary Statistics

Sample period	1947:Q2-2012:Q1
$\bar{x}_{LT}$	1.3
S_{LR}	4.0

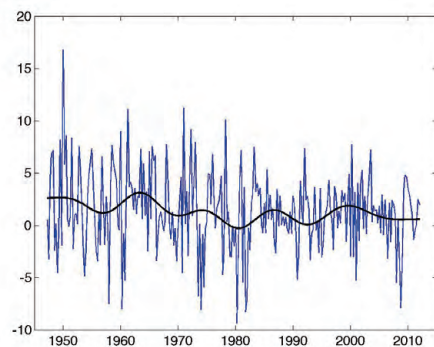
(ii) Cosine Transformations



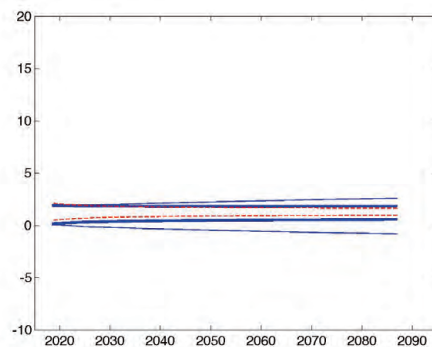
(iii) Low-frequency $I(d)$ log likelihood



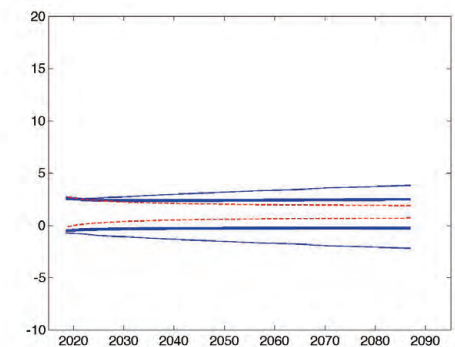
(iv) Data and low-frequency components



(v) 67% Prediction sets

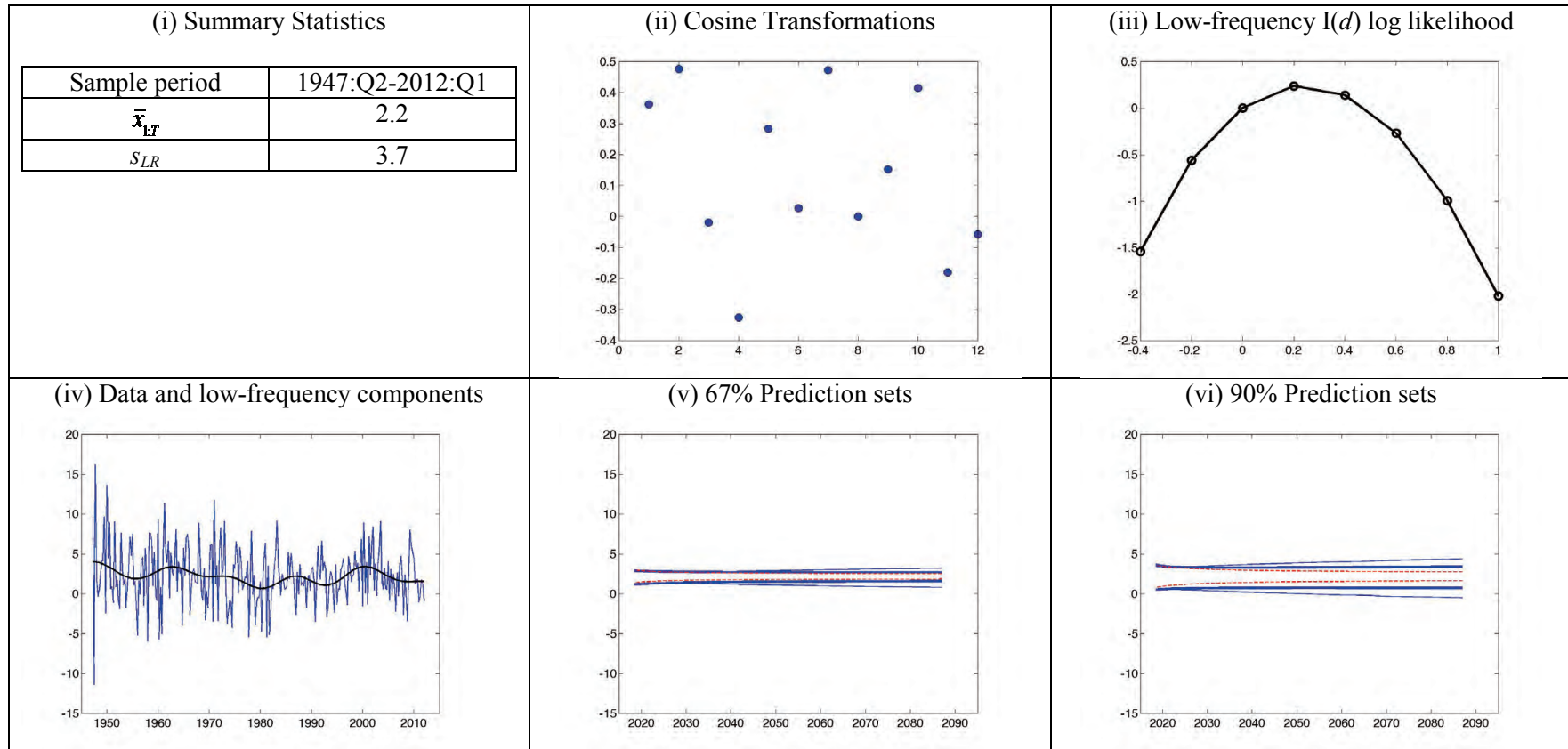


(vi) 90% Prediction sets



See notes to A.1.

A.4 Growth rate of labor productivity



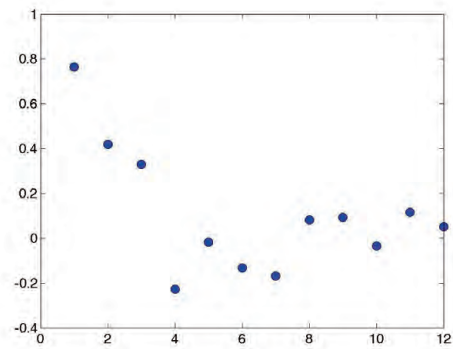
See notes to A.1.

A.5 Growth rate of population

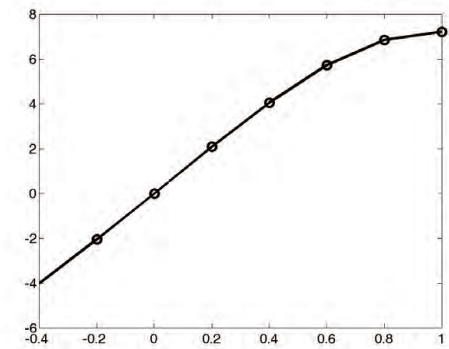
(i) Summary Statistics

Sample period	1947:Q2-2012:Q1
$\bar{x}_{LT}$	1.2
S_{LR}	1.5

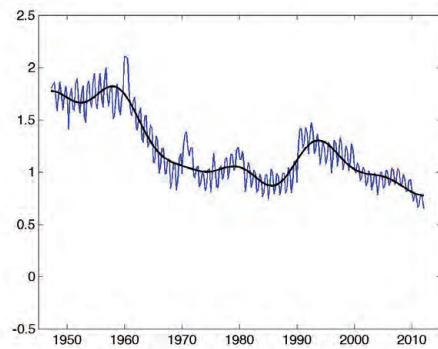
(ii) Cosine Transformations



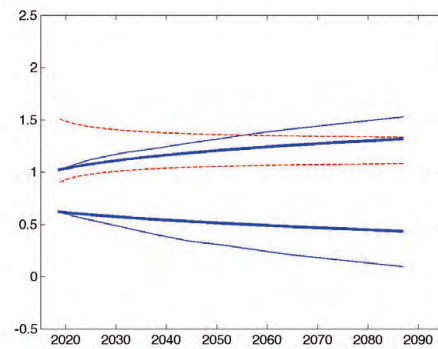
(iii) Low-frequency I(d) log likelihood



(iv) Data and low-frequency components

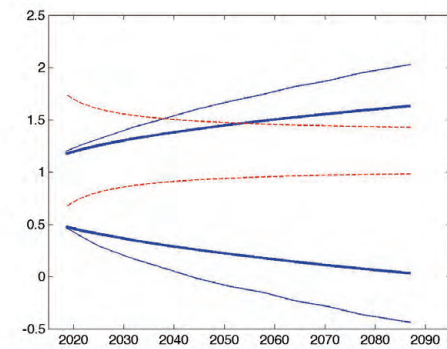


(v) 67%



Prediction sets

(vi) 90% Prediction sets



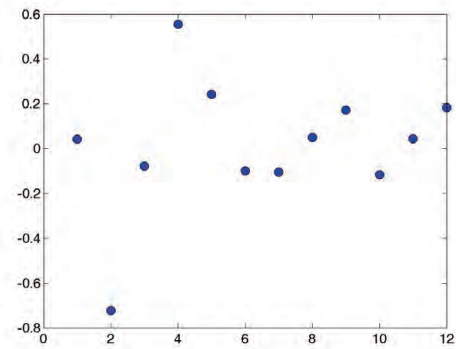
See notes to A.1.

A.6 Inflation (PCE)

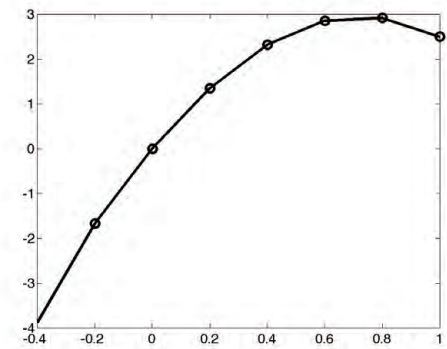
(i) Summary Statistics

Sample period	1947:Q2-2012:Q1
$\bar{x}_{LT}$	3.3
S_{LR}	9.0

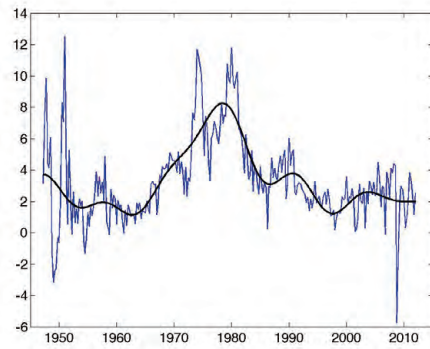
(ii) Cosine Transformations



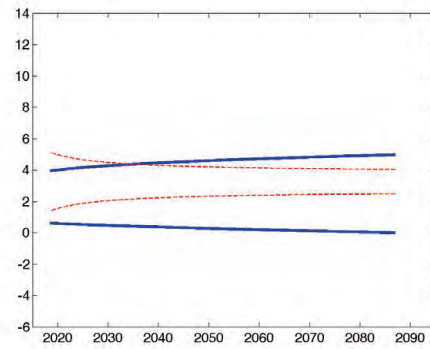
(iii) Low-frequency $I(d)$ log likelihood



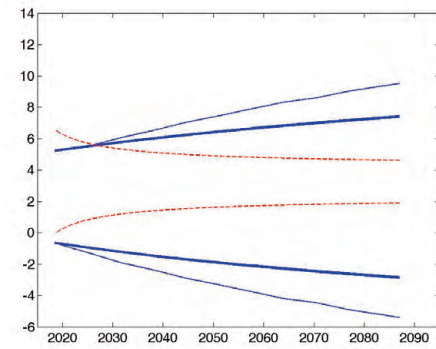
(iv) Data and low-frequency components



(v) 67% Prediction sets



(vi) 90% Prediction sets



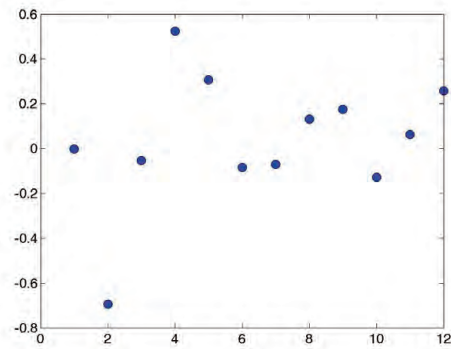
See notes to A.1.

A.7 Inflation (CPI)

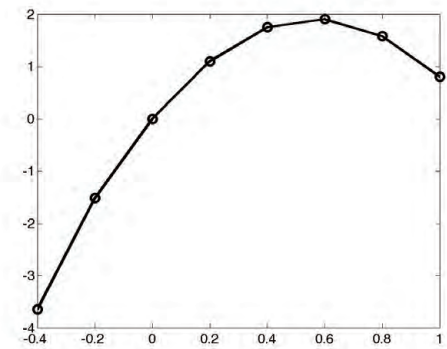
(i) Summary Statistics

Sample period	1947:Q2-2012:Q1
$\bar{x}_{LT}$	3.6
S_{LR}	10.4

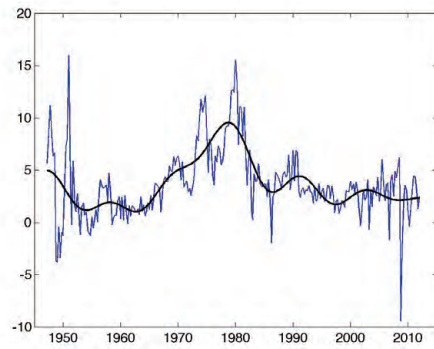
(ii) Cosine Transformations



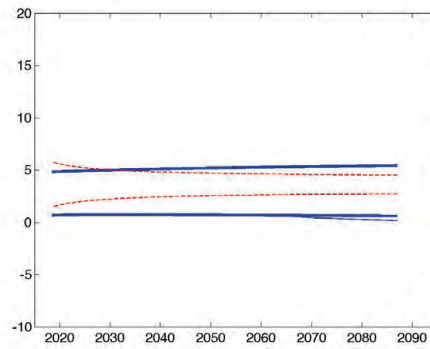
(iii) Low-frequency $I(d)$ log likelihood



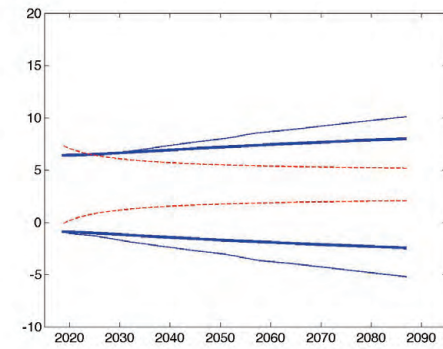
(iv) Data and low-frequency components



(v) 67% Prediction sets



(vi) 90% Prediction sets



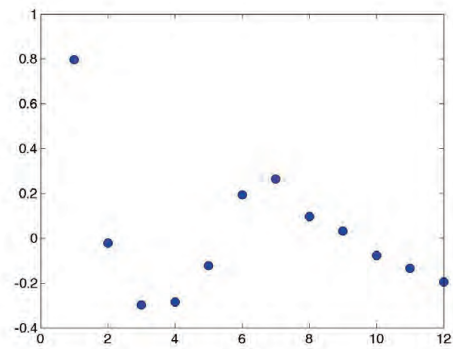
See notes to A.1.

A.8 Inflation (CPI, Japan)

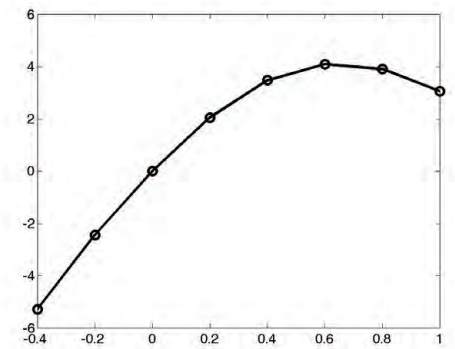
(i) Summary Statistics

Sample period	1947:Q2-2012:Q1
$\bar{x}_{LT}$	3.2
S_{LR}	15.0

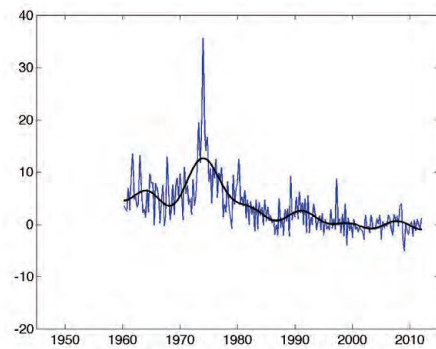
(ii) Cosine Transformations



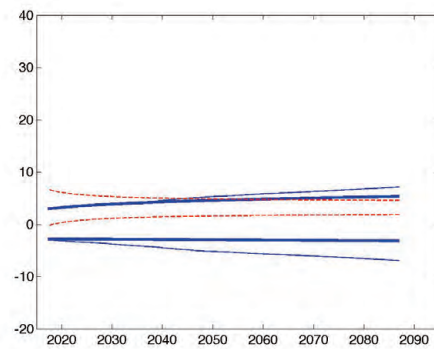
(iii) Low-frequency $I(d)$ log likelihood



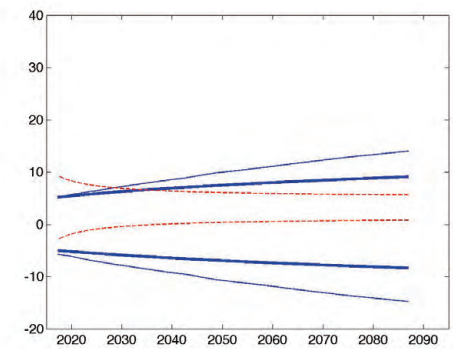
(iv) Data and low-frequency components



(v) 67% Prediction sets



(vi) 90% Prediction sets



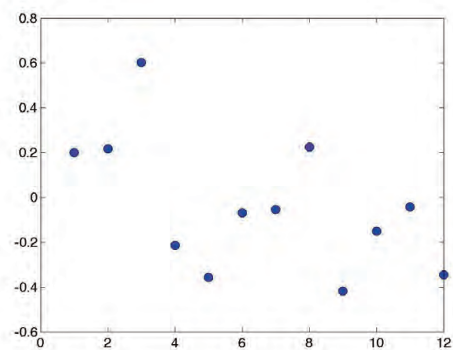
See notes to A.1.

A.9 Stock Returns

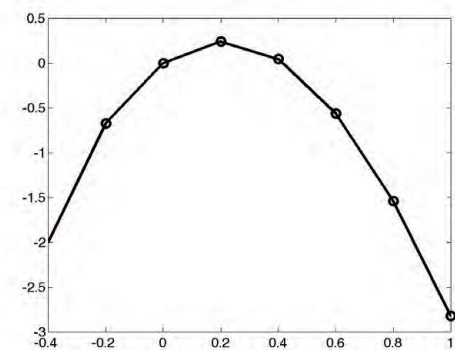
(i) Summary Statistics

Sample period	1947:Q2-2012:Q1
$\bar{x}_{LT}$	6.7
S_{LR}	30.1

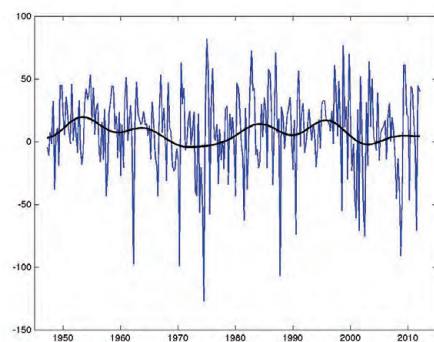
(ii) Cosine Transformations



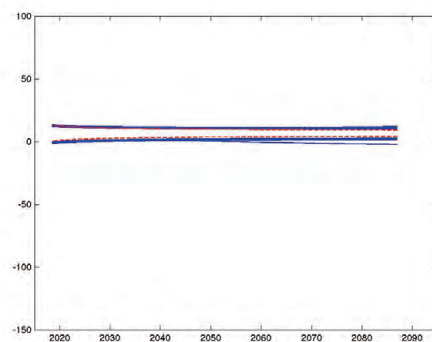
(iii) Low-frequency $I(d)$ log likelihood



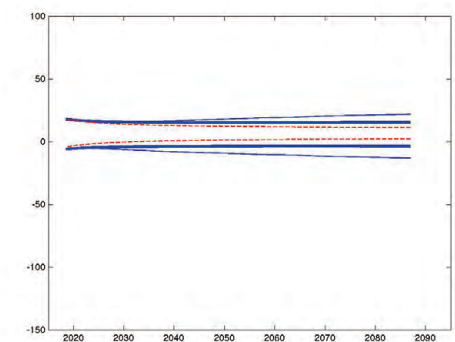
(iv) Data and low-frequency components



(v) 67% Prediction sets

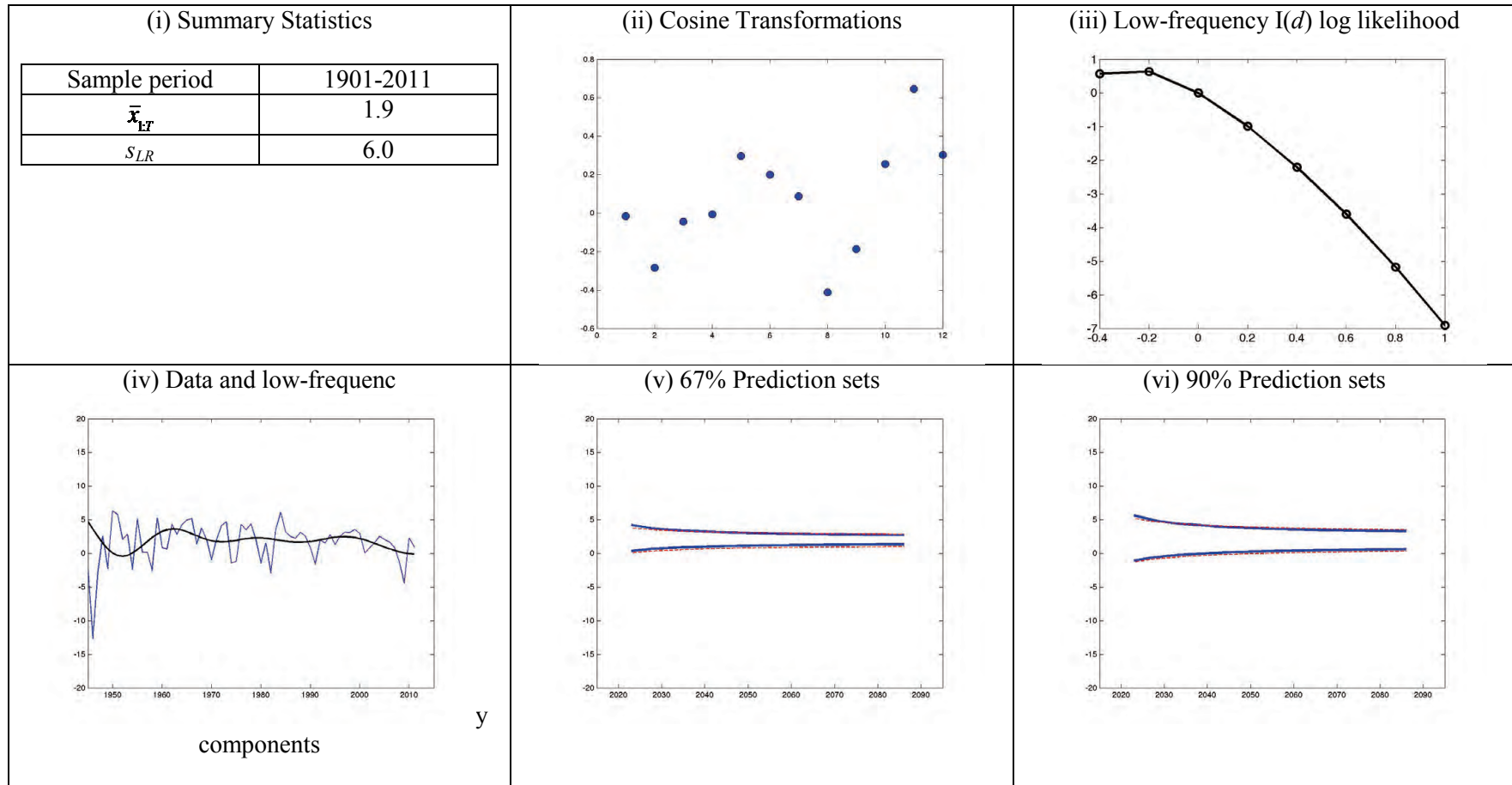


(vi) 90% Prediction sets



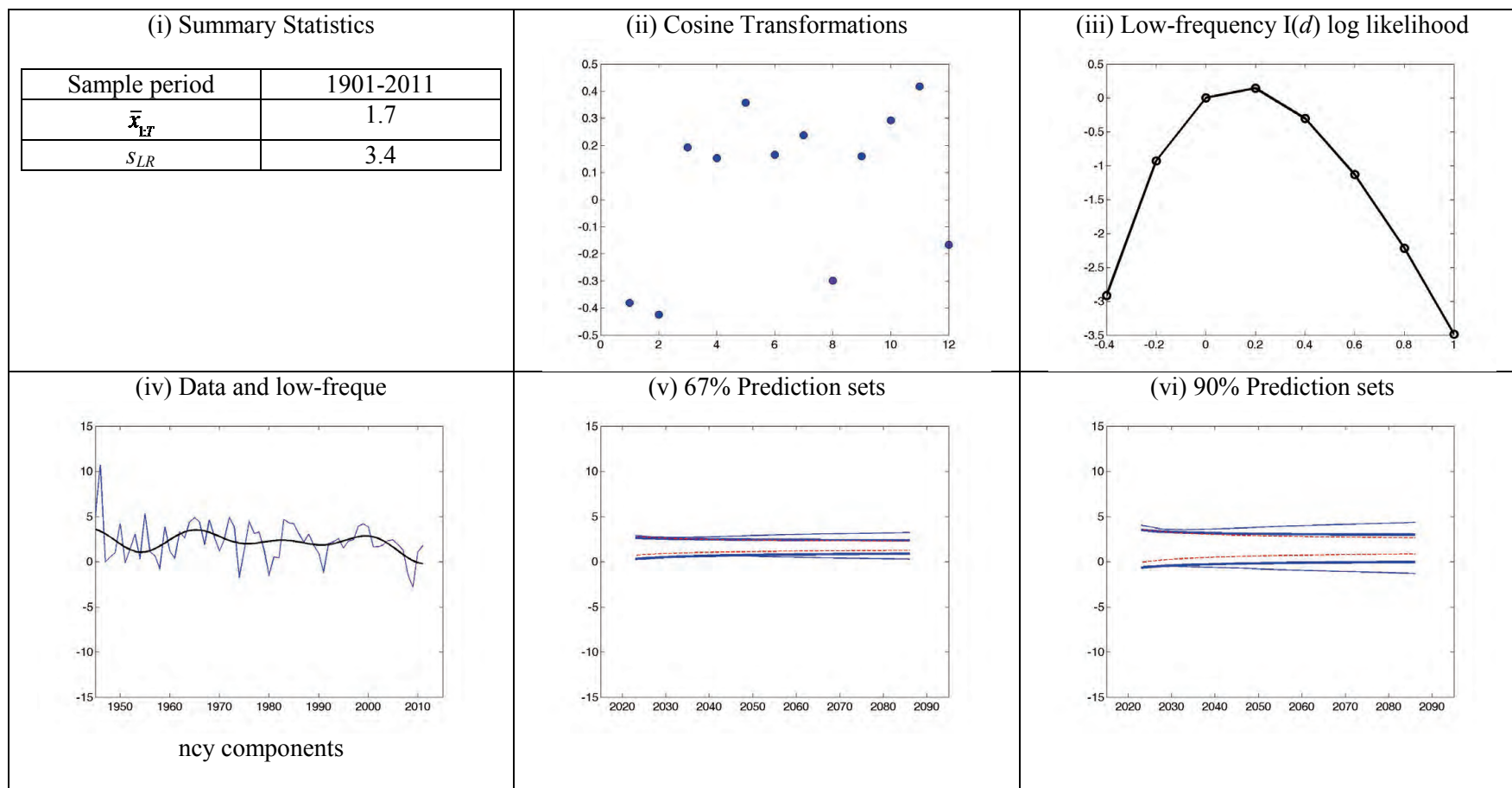
See notes to A.1.

A.10 Growth rate of real per-capita GDP (annual data)



See notes to A.1.

A.11 Growth rate of real per-capita consumption expenditures (annual data)



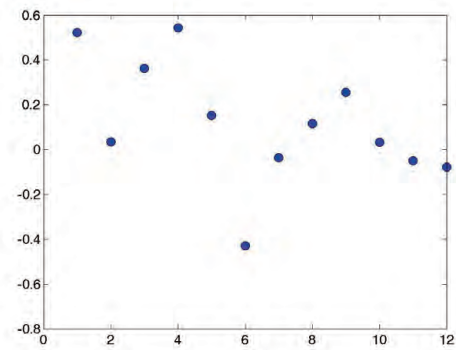
See notes to A.1.

A.12 Growth rate of population (annual data)

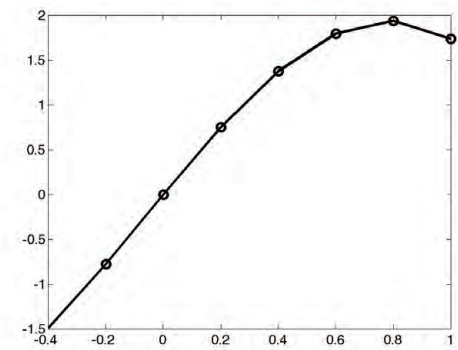
(i) Summary Statistics

Sample period	1901-2011
$\bar{x}_{LT}$	1.3
S_{LR}	1.1

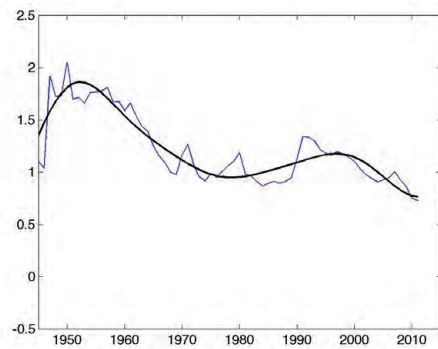
(ii) Cosine Transformations



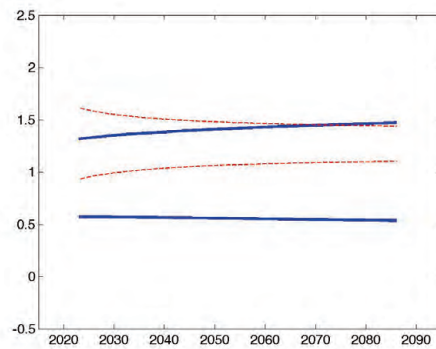
(iii) Low-frequency $I(d)$ log likelihood



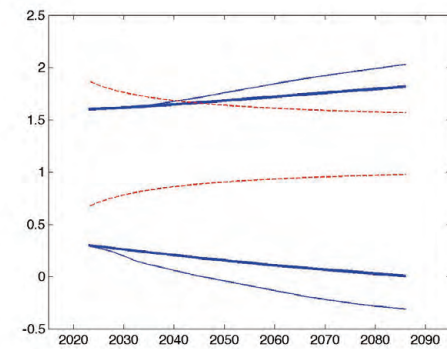
(iv) Data and low-frequency components



(v) 67% Prediction sets

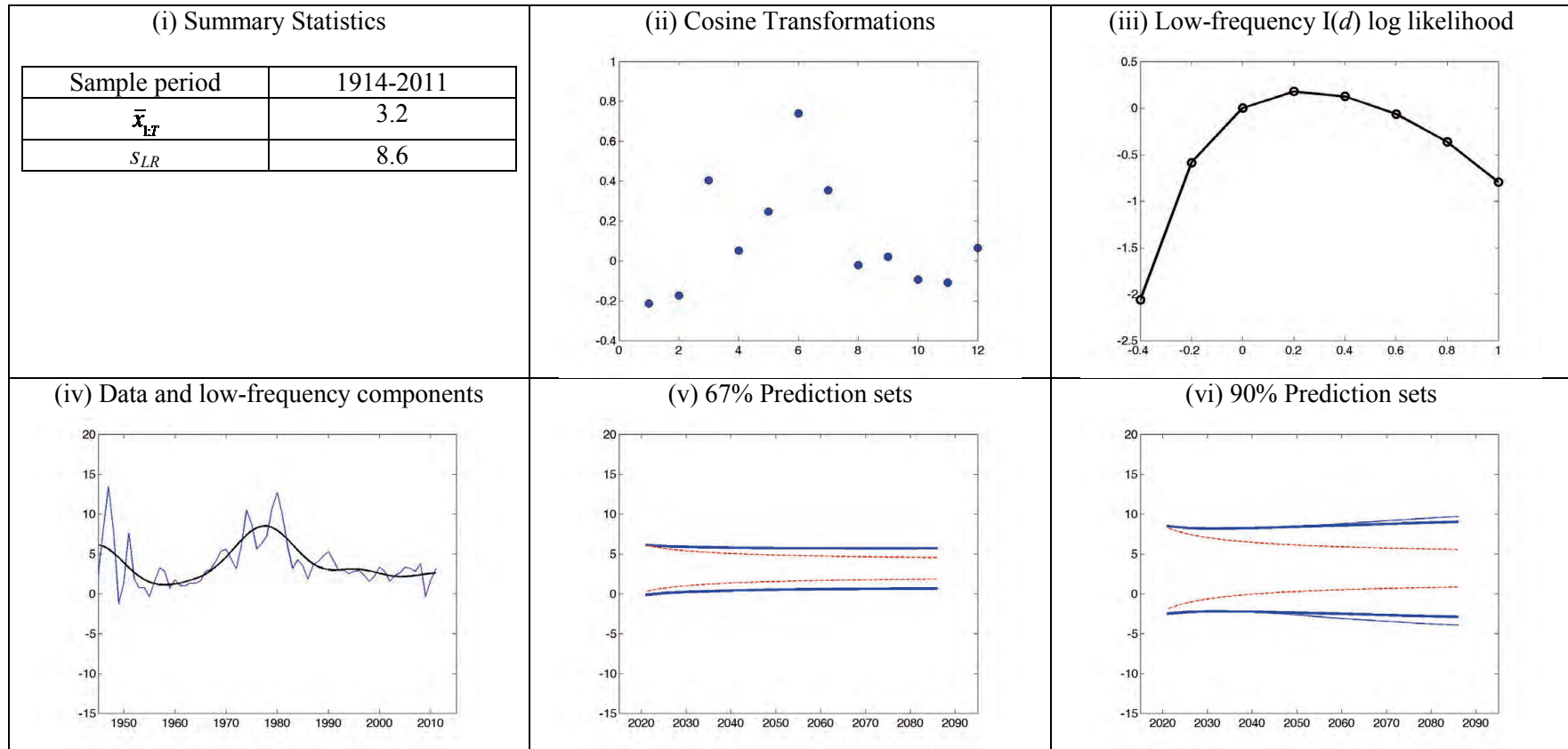


(vi) 90% Prediction sets



See notes to A.1.

A.13 Inflation (CPI, annual data)



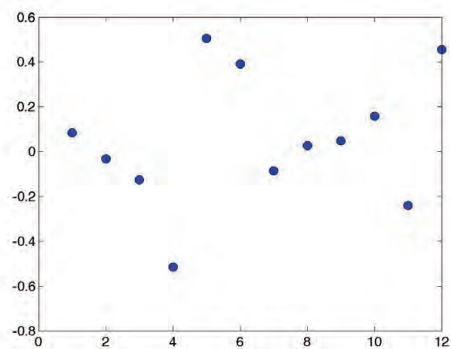
See notes to A.1.

A.14 Stock Returns

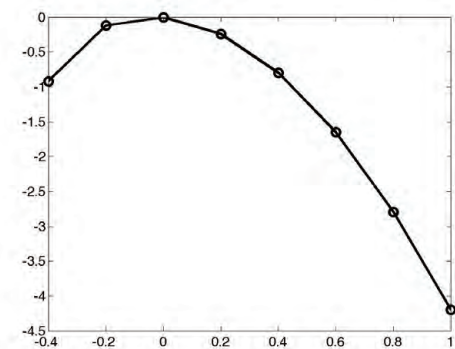
(i) Summary Statistics

Sample period	1926:Q2-2012:Q1
$\bar{x}_{LT}$	6.4
S_{LR}	29.3

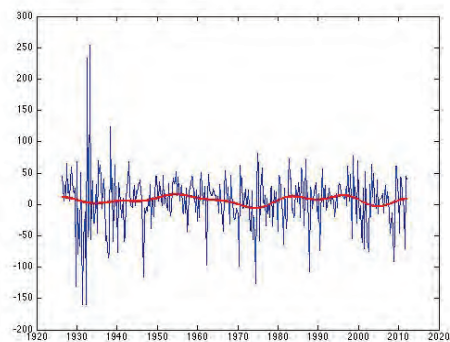
(ii) Cosine Transformations



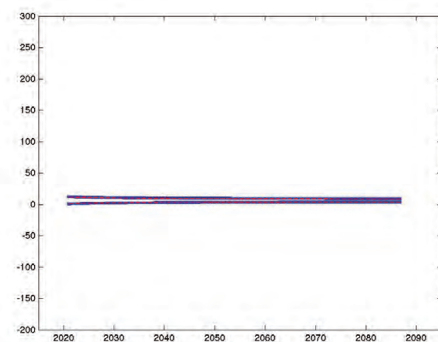
(iii) Low-frequency $I(d)$ log likelihood



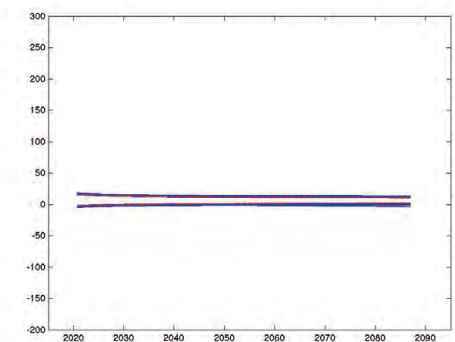
(iv) Data and low-frequency components



(v) 67% Prediction sets



(vi) 90% Prediction sets



See notes to A.1.