

Problem Set 2

Econ 6376 — Time Series Econometrics

Grading

Component	Points
Report craft	20
Code and reproducibility	20
Question 1	35
Question 2	40
Question 3	50
Total	165

Part	Question 1	Question 2	Question 3
(a)	6	12	12
(b)	10	8	12
(c)	7	6	8
(d)	8	10	18
(e)	4	4	—

Question 1 — Stationarity: the DGP and the sample (*Modules 1–3*)

Consider the two processes

$$y_t = 2.5 + 0.80y_{t-1} - 0.43y_{t-2} + \epsilon_{y,t} + 0.7\epsilon_{y,t-1},$$

$$z_t = 3.5 + 0.125t + 0.8z_{t-1} + \epsilon_{z,t} - 0.5\epsilon_{z,t-1}.$$

The innovation sequences are mutually independent, with $\epsilon_{y,t} \stackrel{iid}{\sim} N(0,1)$ and $\epsilon_{z,t} \stackrel{iid}{\sim} N(0,1)$.

- (a) For each process, determine whether it is weakly stationary about a constant mean, stationary about a nonconstant deterministic linear trend, or neither. **Answer from the DGP, without using statistical tests.**
- (b) Set the seed once using `set.seed(11051985)`. Simulate 100 retained observations from each process using a 20-period burn-in. For this part, analyze only the first 50 retained observations of each series. Treat the generating equations as unknown when interpreting the simulated data. Assess whether the evidence supports stationarity about a constant mean, trend stationarity, or a stochastic unit root. Provide all evidence to support your assertion.
- (c) Repeat the analysis using all 100 retained observations from each process. Use the same diagnostic procedure and report whether your assessment changes.
- (d) Compare your theoretical classification with your empirical assessments at $T = 50$ and $T = 100$. Where do they agree or disagree? Explain how sampling variation, test power, and test specification could affect your findings. Does increasing the sample size guarantee that your empirical assessment will match the theoretical classification? Agreement across all three settings is an equally valid outcome when supported by the evidence.
- (e) Suppose you find that Δz_t is stationary. Does this establish that z_t contains a unit root? Explain.

Question 2 — Reading ACF and PACF structure (*Module 3*)

Consider the two processes

$$x_t = 0.4x_{t-1} + \epsilon_t,$$

$$w_t = \eta_t + 0.5\eta_{t-1}.$$

For each realization, the innovations are iid $N(0,1)$, with mutually independent innovation sequences across processes and realizations.

- (a) Derive the population autocorrelations ρ_1 , ρ_2 , and ρ_3 for each process. Describe the population ACF and PACF patterns: which cuts off, at what lag, and which tails off? Explain what the lag-one autocorrelation alone can and cannot tell you about which of these processes generated a series. A full derivation of the MA(1) PACF is not required.
- (b) Write R code to simulate **two independent realizations of each process**, each with 160 retained observations and a 20-period burn-in. For each of the four realizations, set all required presample values and innovations to zero, simulate 180 observations, and discard the first 20. For each retained realization, use its first 80 observations as the $T = 80$ sample and all 160 observations as the $T = 160$ sample. The two sample sizes must therefore be nested within each realization.
- (c) Plot the sample ACF and PACF at lags 1 through 20 for each process, realization, and sample size. This produces eight ACF/PACF pairs (16 panels). Label the process, realization, and sample size clearly; use comparable axes and conventional approximate 95% pointwise reference bands. Choose a readable layout. For example, use one figure per process, with four rows for the two realizations at the two sample sizes and columns for the ACF and PACF.
- (d) Compare the sample patterns with your theoretical results in part (a). At each sample size, how do the two independent realizations of the same process differ? Within each realization, what changes when you increase T from 80 to 160? Treat the DGP as unknown when interpreting each ACF/PACF pair: identify a plausible model class and explain any uncertainty. Report the realizations generated with the specified seed, including ambiguous patterns; your classifications need not change across samples or recover the true model.
- (e) How should you interpret an isolated sample autocorrelation or partial autocorrelation outside the reference bands? Explain why counting the number of spikes outside the bands does not, by itself, identify the model order. Does doubling the sample size guarantee a more accurate model classification for a particular realization? Explain.

Question 3 — Model selection with California unemployment (Modules 4–5)

Download and import the **California unemployment rate**, FRED series [CAUR](#), for **January 2010 through December 2019**, inclusive. This U.S. Bureau of Labor Statistics series is monthly, measured in percent, and seasonally adjusted. Verify that your sample contains 120 consecutive monthly observations with the correct start and end dates and no missing values. Report your retrieval date and retain the raw download so that your analysis can be reproduced.

Your goal is to select a parsimonious model that adequately describes the series' dependence. A smaller final model or a unique preferred model is not guaranteed.

- (a) Explain the general information-criterion structure

$$\text{IC} = -2\hat{\ell} + \text{penalty}(k, n).$$

Define the maximized log-likelihood $\hat{\ell}$, parameter count k , and effective sample size n . Throughout this question, count all estimated coefficients, including any deterministic terms, **plus the innovation variance** in k .

Choose **AIC, BIC, or HQIC before fitting alternative models**. State its formula and defend your choice by comparing the criteria's penalty strengths and the tradeoffs relevant to your modeling objective. Do not choose a criterion simply because its numerical value is smaller than another criterion's value.

HQIC reference: The Hannan–Quinn information criterion is

$$\text{HQIC} = -2\hat{\ell} + 2k \log \log n,$$

where the logarithms are natural. For sample sizes near 120, its penalty per parameter lies between AIC's 2 and BIC's $\log n$. This provides an intermediate penalty strength; it does not guarantee correct model selection in a particular sample.

- (b) Plot the series and assess its stationarity using the visual evidence and appropriately specified ADF and KPSS tests. Explain your choices of deterministic terms and lags. Justify a working decision to retain levels, detrend, or difference, and discuss any conflicting evidence. The series is already seasonally adjusted: monthly frequency alone does not justify seasonal differencing.
- (c) Examine the sample ACF and PACF of your chosen representation. Propose and justify an initial ARMA or ARIMA specification. Explain both the patterns that support your choice and any uncertainty about the orders.
- (d) Fit your initial model and **two or three justified alternatives** by maximum likelihood. Use your chosen information criterion together with residual time plots, residual ACFs, and Ljung–Box tests to guide and assess revisions. Explain the reason for each revision. Include an attempt to simplify the ARMA structure where possible; you may retain a term if its removal worsens the evidence, or add terms if the diagnostics warrant them. If your initial model has no AR or MA terms to remove, explain why no such reduction is available. Provide a compact comparison table showing each model's specification and any supplementary information (e.g., information criterion).