

Introduction to Time Series Analysis. Lecture 10.

Peter Bartlett

Last lecture: Forecasting.

1. Partial autocorrelation function.
2. Recursive methods: Durbin-Levinson.

Introduction to Time Series Analysis. Lecture 10.

1. Review: Forecasting.
2. The innovations representation.
3. Recursive method: Innovations algorithm.
4. Example: Innovations algorithm for forecasting an MA(1)

Review: One-step-ahead linear prediction

$$X_{n+1}^n = \phi_{n1}X_n + \phi_{n2}X_{n-1} + \cdots + \phi_{nn}X_1$$

$$\Gamma_n \phi_n = \gamma_n,$$

$$P_{n+1}^n = E \left(X_{n+1} - X_{n+1}^n \right)^2 = \gamma(0) - \gamma_n' \Gamma_n^{-1} \gamma_n,$$

$$\Gamma_n = \begin{bmatrix} \gamma(0) & \gamma(1) & \cdots & \gamma(n-1) \\ \gamma(1) & \gamma(0) & & \gamma(n-2) \\ \vdots & & \ddots & \vdots \\ \gamma(n-1) & \gamma(n-2) & \cdots & \gamma(0) \end{bmatrix},$$

$$\phi_n = (\phi_{n1}, \phi_{n2}, \dots, \phi_{nn})', \quad \gamma_n = (\gamma(1), \gamma(2), \dots, \gamma(n))'.$$

Review: Durbin-Levinson

$$\phi_0 = 0,$$

$$\phi_{00} = 0;$$

$$\phi_1 = \phi_{11},$$

$$\phi_{11} = \frac{\gamma(1)}{\gamma(0)};$$

$$\phi_n = \begin{pmatrix} \phi_{n-1} - \phi_{nn}\tilde{\phi}_{n-1} \\ \phi_{nn} \end{pmatrix}, \quad \phi_{nn} = \frac{\gamma(n) - \phi'_{n-1}\tilde{\gamma}_{n-1}}{\gamma(0) - \phi'_{n-1}\gamma_{n-1}}.$$

$$\phi_n = (\phi_{n1}, \dots, \phi_{nn})'$$

$$\tilde{\phi}_n = (\phi_{nn}, \dots, \phi_{n1})',$$

$$\gamma_n = (\gamma(1), \dots, \gamma(n))'$$

$$\tilde{\gamma}_n = (\gamma(n), \dots, \gamma(1))'.$$

Durbin-Levinson: Evolution of mean square error

$$\begin{aligned} P_{n+1}^n &= \gamma(0) - \phi_n' \gamma_n \\ &= \gamma(0) - \begin{pmatrix} \phi_{n-1} - \phi_{nn} \tilde{\phi}_{n-1}' \\ \phi_{nn} \end{pmatrix}' \begin{pmatrix} \gamma_{n-1} \\ \gamma(n) \end{pmatrix} \\ &= P_n^{n-1} - \phi_{nn} \left(\gamma(n) - \tilde{\phi}_{n-1}' \gamma_{n-1} \right) \\ &= P_n^{n-1} - \phi_{nn}^2 \left(\gamma(0) - \phi_{n-1}' \gamma_{n-1} \right) \quad (\text{From expression for } \phi_{nn}) \\ &= P_n^{n-1} (1 - \phi_{nn}^2) . \end{aligned}$$

i.e., variance reduces by a factor $1 - \phi_{nn}^2$.

Introduction to Time Series Analysis. Lecture 10.

1. Review: Forecasting.
2. The innovations representation.
3. Recursive method: Innovations algorithm.
4. Example: Innovations algorithm for forecasting an MA(1)

The innovations representation

Instead of writing the best linear predictor as

$$X_{n+1}^n = \phi_{n1}X_n + \phi_{n2}X_{n-1} + \cdots + \phi_{nn}X_1,$$

we can write

$$X_{n+1}^n = \theta_{n1} \underbrace{(X_n - X_n^{n-1})}_{\text{innovation}} + \theta_{n2} (X_{n-1} - X_{n-1}^{n-2}) + \cdots + \theta_{nn} (X_1 - X_1^0).$$

This is still linear in $X_1, \dots, X_n$.

The innovations are uncorrelated:

$$\text{Cov}(X_j - X_j^{j-1}, X_i - X_i^{i-1}) = 0 \text{ for } i \neq j.$$

Comparing representations: $U_n = X_n - X_n^{n-1}$ versus X_n

$\{U_n\}$ form a *decorrelated* representation for the $\{X_n\}$:

$$\begin{pmatrix} U_1 \\ U_2 \\ \vdots \\ U_n \end{pmatrix} = \begin{pmatrix} 1 & 0 & \cdots & 0 \\ -\phi_{11} & 1 & & 0 \\ \vdots & & \ddots & \\ -\phi_{n-1,n-1} & -\phi_{n-1,n-2} & \cdots & 1 \end{pmatrix} \begin{pmatrix} X_1 \\ X_2 \\ \vdots \\ X_n \end{pmatrix}$$

Comparing representations: $U_n = X_n - X_n^{n-1}$ versus X_n

$$\begin{pmatrix} X_1^0 \\ X_2^1 \\ \vdots \\ X_n^{n-1} \end{pmatrix} = \begin{pmatrix} 0 & 0 & \cdots & 0 \\ \theta_{11} & 0 & & 0 \\ \vdots & & \ddots & \\ \theta_{n-1,n-1} & \theta_{n-1,n-2} & \cdots & 0 \end{pmatrix} \begin{pmatrix} U_1 \\ U_2 \\ \vdots \\ U_n \end{pmatrix}$$

Introduction to Time Series Analysis. Lecture 10.

1. Review: Forecasting.
2. The innovations representation.
3. Recursive method: Innovations algorithm.
4. Example: Innovations algorithm for forecasting an MA(1)

Innovations Algorithm

$$X_1^0 = 0, \quad X_{n+1}^n = \sum_{i=1}^n \theta_{ni} (X_{n+1-i} - X_{n+1-i}^{n-i}).$$

$$\theta_{n,n-i} = \frac{1}{P_{i+1}^i} \left(\gamma(n-i) - \sum_{j=0}^{i-1} \theta_{i,i-j} \theta_{n,n-j} P_{j+1}^j \right).$$

$$P_1^0 = \gamma(0) \quad P_{n+1}^n = \gamma(0) - \sum_{i=0}^{n-1} \theta_{n,n-i}^2 P_{i+1}^i.$$

Innovations Algorithm: Example

$$\theta_{n,n-i} = \frac{1}{P_{i+1}^i} \left(\gamma(n-i) - \sum_{j=0}^{i-1} \theta_{i,i-j} \theta_{n,n-j} P_{j+1}^j \right).$$

$$P_1^0 = \gamma(0) \quad P_{n+1}^n = \gamma(0) - \sum_{i=0}^{n-1} \theta_{n,n-i}^2 P_{i+1}^i.$$

$$\theta_{1,1} = \gamma(1)/P_1^0, \quad P_2^1 = \gamma(0) - \theta_{1,1}^2 P_1^0$$

$$\theta_{2,2} = \gamma(2)/P_1^0, \quad \theta_{2,1} = (\gamma(1) - \theta_{1,1} \theta_{2,2} P_1^0) / P_2^1,$$

$$P_3^2 = \gamma(0) - (\theta_{2,2}^2 P_1^0 + \theta_{2,1}^2 P_2^1)$$

$$\theta_{3,3}, \quad \theta_{3,2}, \quad \theta_{3,1}, \quad P_4^3, \dots$$

Predicting h steps ahead using innovations

The innovations representation for the one-step-ahead forecast is

$$P(X_{n+1}|X_1, \dots, X_n) = \sum_{i=1}^n \theta_{ni} (X_{n+1-i} - X_{n+1-i}^{n-i}),$$

What is the innovations representation for $P(X_{n+h}|X_1, \dots, X_n)$?

It is $P(X_{n+h}|X_1, \dots, X_{n+h-1})$, but with the unobserved innovations (from $n+1$ to $n+h-1$) set to zero.

Predicting h steps ahead using innovations

What is the innovations representation for $P(X_{n+h}|X_1, \dots, X_n)$?

Fact: If $h \geq 1$ and $1 \leq i \leq n$, we have

$$\text{Cov}(X_{n+h} - P(X_{n+h}|X_1, \dots, X_{n+h-1}), X_i) = 0.$$

Thus, $P(X_{n+h} - P(X_{n+h}|X_1, \dots, X_{n+h-1})|X_1, \dots, X_n) = 0$.

That is, the best prediction of X_{n+h} is the

best prediction of the one-step-ahead forecast of X_{n+h} .

Fact: The best prediction of $X_{n+1} - X_{n+1}^n$ given only $X_1, \dots, X_n$ is 0.

Similarly for $n+2, \dots, n+h-1$.

Predicting h steps ahead using innovations

Innovations representation:

$$P(X_{n+h}|X_1, \dots, X_n) = \sum_{i=1}^n \theta_{n+h-1, h-1+i} (X_{n+1-i} - X_{n+1-i}^{n-i})$$

Predicting h steps ahead using innovations (Details)

$$\begin{aligned} & P(X_{n+h} | X_1, \dots, X_n) \\ &= P(P(X_{n+h} | X_1, \dots, X_{n+h-1}) | X_1, \dots, X_n) \\ &= P\left(\sum_{i=1}^{n+h-1} \theta_{n+h-1,i} (X_{n+h-i} - X_{n+h-i}^{n+h-i-1}) | X_1, \dots, X_n\right) \\ &= \sum_{i=1}^{n+h-1} \theta_{n+h-1,i} P((X_{n+h-i} - X_{n+h-i}^{n+h-i-1}) | X_1, \dots, X_n) \\ &= \sum_{i=h}^{n+h-1} \theta_{n+h-1,i} P((X_{n+h-i} - X_{n+h-i}^{n+h-i-1}) | X_1, \dots, X_n) \\ &= \sum_{i=h}^{n+h-1} \theta_{n+h-1,i} (X_{n+h-i} - X_{n+h-i}^{n+h-i-1}) \end{aligned}$$

Predicting h steps ahead using innovations (Details)

$$P(X_{n+1}|X_1, \dots, X_n) = \sum_{i=1}^n \theta_{ni} (X_{n+1-i} - X_{n+1-i}^{n-i})$$

$$\begin{aligned} P(X_{n+h}|X_1, \dots, X_n) &= \sum_{j=h}^{n+h-1} \theta_{n+h-1,j} (X_{n+h-j} - X_{n+h-j}^{n+h-j-1}) \\ &= \sum_{i=1}^n \theta_{n+h-1,h-1+i} (X_{n+1-i} - X_{n+1-i}^{n-i}) \end{aligned}$$

$$(j = i + h - 1)$$

Mean squared error of h -step-ahead forecasts

From orthogonality of the predictors and the error,

$$E((X_{n+h} - P(X_{n+h}|X_1, \dots, X_n)) P(X_{n+h}|X_1, \dots, X_n)) = 0.$$

That is, $E(X_{n+h} P(X_{n+h}|X_1, \dots, X_n)) = E(P(X_{n+h}|X_1, \dots, X_n)^2)$.

Hence, we can express the mean squared error as

$$\begin{aligned} P_{n+h}^n &= E(X_{n+h} - P(X_{n+h}|X_1, \dots, X_n))^2 \\ &= \gamma(0) + E(P(X_{n+h}|X_1, \dots, X_n))^2 \\ &\quad - 2E(X_{n+h} P(X_{n+h}|X_1, \dots, X_n)) \\ &= \gamma(0) - E(P(X_{n+h}|X_1, \dots, X_n))^2. \end{aligned}$$

Mean squared error of h -step-ahead forecasts

But the innovations are uncorrelated, so

$$\begin{aligned} P_{n+h}^n &= \gamma(0) - \mathbb{E} \left(P(X_{n+h} | X_1, \dots, X_n) \right)^2 \\ &= \gamma(0) - \mathbb{E} \left(\sum_{j=h}^{n+h-1} \theta_{n+h-1,j} \left(X_{n+h-j} - X_{n+h-j}^{n+h-j-1} \right) \right)^2 \\ &= \gamma(0) - \sum_{j=h}^{n+h-1} \theta_{n+h-1,j}^2 \mathbb{E} \left(X_{n+h-j} - X_{n+h-j}^{n+h-j-1} \right)^2 \\ &= \gamma(0) - \sum_{j=h}^{n+h-1} \theta_{n+h-1,j}^2 P_{n+h-j}^{n+h-j-1}. \end{aligned}$$

Introduction to Time Series Analysis. Lecture 10.

1. Review: Forecasting.
2. The innovations representation.
3. Recursive method: Innovations algorithm.
4. Example: Innovations algorithm for forecasting an MA(1)

Example: Innovations algorithm for forecasting an MA(1)

Suppose that we have an MA(1) process $\{X_t\}$ satisfying

$$X_t = W_t + \theta_1 W_{t-1}.$$

Given $X_1, X_2, \dots, X_n$, we wish to compute the best linear forecast of X_{n+1} , using the innovations representation,

$$X_1^0 = 0, \quad X_{n+1}^n = \sum_{i=1}^n \theta_{ni} (X_{n+1-i} - X_{n+1-i}^{n-i}).$$

Example: Innovations algorithm for forecasting an MA(1)

An aside: The linear predictions are in the form

$$X_{n+1}^n = \sum_{i=1}^n \theta_{ni} Z_{n+1-i}$$

for uncorrelated, zero mean random variables Z_i . In particular,

$$X_{n+1} = Z_{n+1} + \sum_{i=1}^n \theta_{ni} Z_{n+1-i},$$

where $Z_{n+1} = X_{n+1} - X_{n+1}^n$ (and all the Z_i are uncorrelated).

This is suggestive of an MA representation.

Why isn't it an MA?

Example: Innovations algorithm for forecasting an MA(1)

$$\theta_{n,n-i} = \frac{1}{P_{i+1}^i} \left(\gamma(n-i) - \sum_{j=0}^{i-1} \theta_{i,i-j} \theta_{n,n-j} P_{j+1}^j \right).$$

$$P_1^0 = \gamma(0) \quad P_{n+1}^n = \gamma(0) - \sum_{i=0}^{n-1} \theta_{n,n-i}^2 P_{i+1}^i.$$

The algorithm computes $P_1^0 = \gamma(0)$, $\theta_{1,1}$ (in terms of $\gamma(1)$);
 P_2^1 , $\theta_{2,2}$ (in terms of $\gamma(2)$), $\theta_{2,1}$; P_3^2 , $\theta_{3,3}$ (in terms of $\gamma(3)$), etc.

Example: Innovations algorithm for forecasting an MA(1)

$$\theta_{n,n-i} = \frac{1}{P_{i+1}^i} \left(\gamma(n-i) - \sum_{j=0}^{i-1} \theta_{i,i-j} \theta_{n,n-j} P_{j+1}^j \right).$$

For an MA(1), $\gamma(0) = \sigma^2(1 + \theta_1^2)$, $\gamma(1) = \theta_1\sigma^2$.

Thus: $\theta_{1,1} = \gamma(1)/P_1^0$;

$\theta_{2,2} = 0$, $\theta_{2,1} = \gamma(1)/P_2^1$;

$\theta_{3,3} = \theta_{3,2} = 0$; $\theta_{3,1} = \gamma(1)/P_3^2$, etc.

Because $\gamma(n-i) \neq 0$ only for $i = n-1$, only $\theta_{n,1} \neq 0$.

Example: Innovations algorithm for forecasting an MA(1)

For the MA(1) process $\{X_t\}$ satisfying

$$X_t = W_t + \theta_1 W_{t-1},$$

the innovations representation of the best linear forecast is

$$X_1^0 = 0, \quad X_{n+1}^n = \theta_{n1} (X_n - X_n^{n-1}).$$

More generally, for an MA(q) process, we have $\theta_{ni} = 0$ for $i > q$.

Example: Innovations algorithm for forecasting an MA(1)

For the MA(1) process $\{X_t\}$,

$$X_1^0 = 0, \quad X_{n+1}^n = \theta_{n1} (X_n - X_n^{n-1}).$$

This is consistent with the observation that

$$X_{n+1} = Z_{n+1} + \sum_{i=1}^n \theta_{ni} Z_{n+1-i},$$

where the uncorrelated Z_i are defined by $Z_t = X_t - X_t^{t-1}$ for $t = 1, \dots, n+1$.

Indeed, as n increases, $P_{n+1}^n \rightarrow \text{Var}(W_t)$ (recall the recursion for P_{n+1}^n), and $\theta_{n1} = \gamma(1)/P_n^{n-1} \rightarrow \theta_1$.

Recall: Forecasting an AR(p)

For the AR(p) process $\{X_t\}$ satisfying

$$X_t = \sum_{i=1}^p \phi_i X_{t-i} + W_t,$$

$$X_1^0 = 0, \quad X_{n+1}^n = \sum_{i=1}^p \phi_i X_{n+1-i}$$

for $n \geq p$. Then

$$X_{n+1} = \sum_{i=1}^p \phi_i X_{n+1-i} + Z_{n+1},$$

where $Z_{n+1} = X_{n+1} - X_{n+1}^n$.

The Durbin-Levinson algorithm is convenient for AR(p) processes.

The innovations algorithm is convenient for MA(q) processes.