

04/03 - Prediction

It is possible to represent a SP in several way, but these representation are not unique. For a given SSP $y(t)$ there may exist infinitely many $W(z)$ and $e(t)$ that reproduce $y(t)$ as output



Case 1: Consider a simple process $y(t) = W(z)e(t)$ with $e(t) \sim WN(0, \lambda^2)$. We can write, for $\alpha \in \mathbb{R}$, that

$$y(t) = \frac{\alpha}{\alpha} W(z)e(t)$$

Define $\tilde{W}(z) = \frac{1}{\alpha} W(z)$ and $\tilde{e}(t) = \alpha e(t)$ that $\tilde{e}(t) \sim WN(0, \alpha^2 \lambda^2)$. So I can write the process as

$$y(t) = \tilde{W}(z)\tilde{e}(t)$$

Case 2: Consider a simple process $y(t) = W(z)e(t)$ with $e(t) \sim WN(0, \lambda^2)$. We can write

$$y(t) = \frac{z^n}{z^n} W(z)e(t)$$

Define $\tilde{W}(z) = z^n W(z)$ and $\tilde{e}(t) = z^{-n} e(t) = e(t-n)$ so that $\tilde{e}(t) \sim WN(0, \lambda^2)$. So I can write the process as

$$y(t) = \tilde{W}(z)\tilde{e}(t)$$

Case 3: Consider a simple process $y(t) = W(z)e(t)$ with $e(t) \sim WN(0, \lambda^2)$. We can write, for $p \in \mathbb{C}$, $|p| < 1$,

$$y(t) = W(z) \frac{z^{-p}}{z^{-p}} e(t)$$

Define $\tilde{W}(z) = W(z) \frac{z^{-p}}{z^{-p}}$ (obtaining a higher degree polynomial - different from the original but same dynamic) and

$\tilde{e}(t) = e(t)$. So I can write the process as

$$y(t) = \tilde{W}(z)\tilde{e}(t)$$

Case 4: Consider a simple process $y(t) = W(z)e(t)$ with $e(t) \sim WN(0, \lambda^2)$.

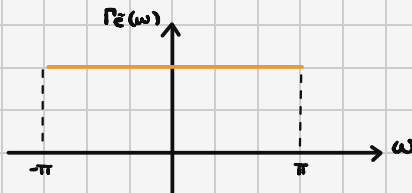
We can write $W(z)$ as $W(z) = W_1(z) \frac{z^{-p}}{z^{-p}}$ so $y(t) = W_1(z) \frac{z^{-p}}{z^{-p}} e(t)$. Define $\tilde{W}(z) = W_1(z) \frac{z^{-p}}{z^{-p}}$ and $\tilde{e}(t) = \frac{z^{-p}}{z^{-p}} e(t)$ (which is still a white noise). So I can write the process as

$$y(t) = \tilde{W}(z)\tilde{e}(t)$$

Proof: $\tilde{e}(t) = \frac{z^{-p}}{z^{-p}} e(t)$ is a white noise

We know that $\tilde{e}(t)$ is computed as output of $G(z)$, so the spectrum of the process $\tilde{e}(t)$ will be defined as

$$\Gamma_{\tilde{e}}(\omega) = \Gamma_e(\omega) |G(e^{j\omega})|^2 = \Gamma_e(\omega) \left| \frac{e^{j\omega} - p}{e^{j\omega} - 1/p} \right|^2$$



$$= \frac{(e^{j\omega} - p)(e^{-j\omega} - \bar{p})}{(e^{j\omega} - 1/p)(e^{-j\omega} - 1/\bar{p})} \lambda^2$$

$$= \frac{1 + p^2 - p(e^{j\omega} + e^{-j\omega})}{1 + 1/p^2 - 1/p(e^{j\omega} + e^{-j\omega})} \lambda^2 = \frac{1 + p^2 - 2p \cos \omega}{1 + 1/p^2 - 2/p \cos \omega} \lambda^2$$

$$= \frac{p^2(1/p^2 + 1 - 1/p(2 \cos \omega))}{1 + 1/p^2 - 2/p \cos \omega} \lambda^2 = p^2 \lambda^2 \quad \forall \omega$$

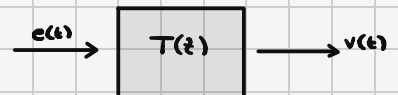
So the shape is again flat, as happened for the white noise. So yes, $\tilde{e}(t)$ is a white noise with scaled variance

$$T = \frac{z^{-p}}{z^{-p}} \cdot \frac{1}{p}$$

is called **ALL-PASS FILTER** (all frequencies are buffered).

$$\Gamma_y(\omega) = |T(e^{j\omega})|^2 \Gamma_e(\omega) = \Gamma_e(\omega)$$

$y(t)$ and $e(t)$ are equivalent (not equal): $e(t) \sim WN(0, \lambda^2)$ $y(t) \sim WN(0, \lambda^2)$



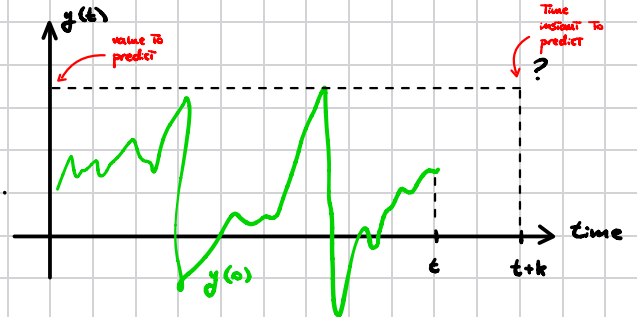
In the following we will use to define any SSP in a unique way. The following result allows us to select the so called **CANONICAL REPRESENTATION**.

THEOREM: Let $y(t)$ be a SSP with $\Gamma_y(w)$ rational. There exists a unique pair $(W(z), e(t))$ if and only if, given $W(z) = C(z)/A(z)$

- 1) $C(z)$ and $A(z)$ are **monic** ($C(z) = 1 + c_1 z^{-1} + \dots$ $A(z) = 1 + a_1 z^{-1} + \dots$)
- 2) $C(z)$ and $A(z)$ have **null relative degree** ($d = m - n \sim$ this means that $d = \# \text{poles} - \# \text{zeros}$)
- 3) $C(z)$ and $A(z)$ are **coprime** (no common factor)
- 4) $C(z)$ and $A(z)$ have roots inside the unit circle

Given a SSP $y(t)$ in canonical form and a specific realization $y(0), \dots, y(t)$, how can we predict the future value of the process at a generic time $t+k$?

$$y(t) = \frac{C(z)}{A(z)} e(t) \quad e(t) \sim WN(\mu, \lambda^2) \quad C(z), A(z), \mu, \lambda^2$$



We denote $\hat{y}(t+k|t)$ the **PREDICTOR** at time $t+k$, given t .
(or $\hat{y}(t|t-1)$ the predictor at time t given a past sample at $t-1$)

Example of predictors:

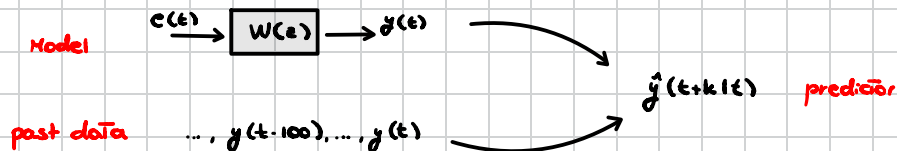
$$1) \quad y(t) = W(z) e(t) = \frac{C(z)}{A(z)} e(t) \quad \Rightarrow \quad m_y = \frac{C(1)}{A(1)} m_e$$

$$2) \quad \hat{y}(t+k|t) = y(t) \quad \text{or} \quad \hat{y}(t+k|t) = \frac{1}{3} (y(t) + y(t-1) + y(t-2)) \quad \text{or} \quad \hat{y}(t+k|t) = \frac{1}{3} (2y(t) + \frac{1}{2}y(t-1) + \frac{1}{2}y(t-2))$$

Under some assumption, they are all good "naive" predictors, but:

- 1) uses only information about the model ($c_1, \dots, c_n, a_1, \dots, a_m$ are all model parameters)
- 2) uses only the data

We would like $\hat{y}(t+k|t)$ to be "optimal" in some sense (so we need an **optimality criteria**) and we would like to exploit all the information we have, so both model and data values



How can we measure the quality of a predictor?

MEAN SQUARE prediction error:

$$J(\hat{y}) = E[(y(t+k) - \hat{y}(t+k|t))^2] = E[\varepsilon(t+k|t)^2]$$

which predictor $\hat{y}(t+k|t)$ minimizes $J(\hat{y})$? The predictor must be a function of both model information and data

$$\hat{y}(t+k|t) = f(\text{model}, \text{data}) \quad \text{HARD TO SOLVE (in the general setting)}$$

Since we work with linear models, we will focus on **linear predictors**

Linear predictors:

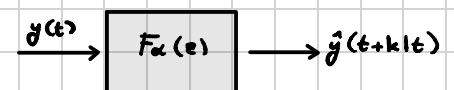
$$\hat{y}(t+k|t) = \alpha_0 y(t) + \alpha_1 y(t-1) + \dots = \sum_{i=0}^{\infty} \alpha_i y(t-i) \quad \alpha_i \in \mathbb{R} \text{ such that } \sum \alpha_i < \infty$$

where

- $\alpha_0, \alpha_1, \dots$ are parameters about the model
- $y(t), y(t-1)$ refers to the available data

Using the backward shift operator, the predictor will be written as

$$\hat{y}(t+k|t) = (\alpha_0 + \alpha_1 z^{-1} + \dots) y(t) = F_\alpha(z) y(t)$$



The goal now is to find $\alpha_0, \alpha_1, \dots$ of the linear predictor $\hat{y}(t+k|t, \alpha) = \sum \alpha_i y(t-i)$ so that the mean square prediction error $J(\alpha) = E[(y(t+k) - \hat{y}(t+k|t, \alpha))^2]$ is minimized.

$$\hat{\alpha} = \underset{\alpha}{\text{argmin}} E[(y(t+k|t) - \hat{y}(t+k|t, \alpha))^2]$$

Solution of The design problem: Consider a simple process $y(t) = W(z)e(t)$ with $e(t) \sim WN(0, \lambda^2)$. We can write

The ARMA $y(t)$ into an MA(∞)

$$y(t) = w_0 e(t) + w_1 e(t-1) + w_2 e(t-2) + \dots = \sum_{i=0}^{+\infty} w_i e(t-i)$$

but at time $t-1$, the process will simply be the time translation of $y(t)$

$$y(t-1) = w_0 e(t-1) + w_1 e(t-2) + \dots = \sum_{i=0}^{+\infty} w_i e(t-i-1)$$

At the same time, the predictor can be written as

$$\hat{y}(t+k|t) = \alpha_0 \sum_{i=0}^{+\infty} w_i e(t-i) + \alpha_1 \sum_{i=0}^{+\infty} w_i e(t-i-1) + \dots$$

So $\hat{y}$ is written as past values of $e(t)$ (past values of white noise)

$$\hat{y}(t+k|t) = \beta_0 e(t) + \beta_1 e(t-1) + \dots = \sum_{i=0}^{+\infty} \beta_i e(t-i)$$

Reformulate the optimization problem with respect to α in terms of a new optimization problem with respect to β .

$$V(\beta) = E[(y(t+k) - \hat{y}(t+k|t, \beta))^2]$$

and $\hat{\beta} = \arg\min_{\beta} V(\beta)$. In this way, it will be easier to find the solution. We can write the prediction at time $t+k$ as

$$y(t+k) = \sum_{i=0}^{+\infty} w_i e(t-i+k) = \underbrace{\sum_{j=0}^{k-1} w_j e(t+k-j)}_{\text{combination of future sample of } e(t)} + \underbrace{\sum_{i=0}^{+\infty} w_{i+k} e(t-i)}_{\text{past / present value of } e(t)}$$

The optimization problem $V(\beta)$ in this way become

$$\begin{aligned} V(\beta) &= E[(y(t+k) - \hat{y}(t+k|t))^2] \\ &= E[(\sum_{j=0}^{k-1} w_j e(t+k-j) + \sum_{i=0}^{+\infty} w_{i+k} e(t-i) - \sum_{i=0}^{+\infty} \beta_i e(t-i))^2] \\ &= E[(\underbrace{\sum_{j=0}^{k-1} w_j e(t+k-j)}_{\text{not function of } \beta \text{ so offset}} + (\sum_{i=0}^{+\infty} (w_{i+k} - \beta_i) e(t-i))^2 + 2(\sum_{j=0}^{k-1} w_j e(t+k-j))(\sum_{i=0}^{+\infty} (w_{i+k} - \beta_i) e(t-i))] \end{aligned}$$

$\rightarrow$ all $e(t)$'s are uncorrelated at different time instant (never coincide in time instant $\rightarrow$ always 0)

The optimal predictor is given by $E[\sum_{i=0}^{+\infty} (w_{i+k} - \beta_i) e(t-i)^2] \geq 0 \Leftrightarrow \beta_i = w_{i+k} \quad \forall i=0, \dots$

optimal linear predictor: From the noise $\rightarrow$ we need $e(t)$ and infinite amount of coefficient (not computable)

$$\hat{y}(t+k|t) = \sum_{i=0}^{+\infty} w_{i+k} e(t-i)$$