

10/03 - The Maximum Likelihood approach

Identification of ARMA/ARMAX models

For an ARMA/ARMAX model we consider a model class defined as:

$$H_\theta : y(t) = \frac{B(z)}{A(z)} u(t-d) + \frac{C(z)}{A(z)} e(t) \quad e(t) \sim \text{WN}(0, \sigma^2)$$

where (written in canonical form):

- $A(z) = 1 - a_1 z^{-1} - \dots - a_m z^{-m}$
- $B(z) = b_0 + b_1 z^{-1} + \dots + b_p z^{-p}$
- $C(z) = 1 + c_1 z^{-1} + \dots + c_n z^{-n}$

and The set of parameter is $\theta = [a_1 \dots a_m \ b_0 \dots b_p \ c_1 \dots c_n]^T$ and The overall size of θ is $n_\theta = m + p + n$
So

$$J_N(\theta) = \frac{1}{N} \sum (y(t) - \hat{y}(t|t-1, \theta))^2$$

The predictor, in ARMAX case is $\hat{y}(t+k|t, \theta) = \frac{B(z)}{C(z)} E(z) u(t+k-d) + \frac{F(z)}{C(z)} y(t, k)$. If we want The 1-step ahead, using The long division, we find That (for $k=1$) we perform only 1 step of long division, and

$$E(z) = 1 \quad F(z)z^{-k} = C(z) - A(z)$$

The final predictor for The ARMA/ARMAX model will be

$C(z)$	$A(z)$
$-A(z)$	$1 \quad E(z)$
$F(z)z^{-1} \quad C(z) - A(z)$	

$$\hat{y}(t|t-1) = \frac{C(z)-A(z)}{C(z)} y(t) + \frac{B(z)}{C(z)} u(t-d)$$

$$\text{prediction error: } \varepsilon(t|t-1, \theta) = y(t) - \hat{y}(t|t-1, \theta) = y(t) - \frac{C(z)-A(z)}{C(z)} y(t) - \frac{B(z)}{C(z)} u(t-d)$$

$$= \left(1 - \frac{C(z)-A(z)}{C(z)}\right) y(t) - \frac{B(z)}{C(z)} u(t-d)$$

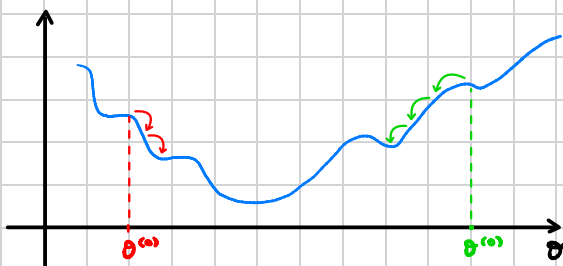
$$= \frac{A(z)}{C(z)} y(t) - \frac{B(z)}{C(z)} u(t-d)$$

θ appears in the denominator of $\varepsilon(t|t-1, \theta)$

Now $\hat{y}(t|t-1, \theta) \neq \Psi^T(t) \theta$ non-linear optimization problem

We need an iterative numerical approach

- we initialize The algorithm with an initial guess $\theta^{(0)}$
- we define an update rule: $\theta^{(i+1)} = f(\theta^{(i)})$. The sequence of The estimates should converge To $\hat{\theta}_N = \arg \min_\theta J_N(\theta)$



We have Two problems:

- 1) Update rule
- 2) multiple minima $\Rightarrow$ many local optimum

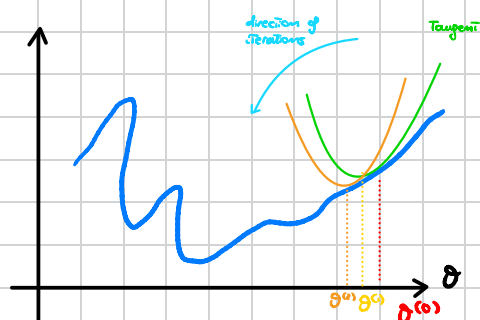
Problem of local optimal

Iterative algorithms are guaranteed To converge To a local minimum, not To The global one. We will follow an heuristic approach

Iterations	$\theta^{(0,1)} \rightarrow \theta^{(1,1)} \rightarrow \theta^{(N,1)} \Rightarrow J_N(\theta^{(N,1)})$
	$\theta^{(0,2)} \rightarrow \theta^{(1,2)} \rightarrow \theta^{(N,2)} \Rightarrow J_N(\theta^{(N,2)})$
	$\vdots$
	$\theta^{(0,M)} \rightarrow \theta^{(1,M)} \rightarrow \theta^{(N,M)} \Rightarrow J_N(\theta^{(N,M)})$

We will Take The one
Giving The lowest cost

How To update $\theta^{(i)}$? We can use the Newton parabola method. Mathematically we define the Tangent paraboloid as (2^o order Taylor expansion)



$$V^{(i)}(\theta) = J_N(\theta^{(i)}) + \left. \frac{\partial}{\partial \theta} J_N(\theta) \right|_{\theta=\theta^{(i)}} (\theta - \theta^{(i)}) + \frac{1}{2} (\theta - \theta^{(i)})^T \left. \frac{\partial^2}{\partial \theta^2} J_N(\theta) \right|_{\theta=\theta^{(i)}} (\theta - \theta^{(i)})$$

The idea is to take $\theta^{(i+1)}$ as the minimizer of $V^{(i)}(\theta)$. So we have to find that

$$\left. \frac{\partial}{\partial \theta} V^{(i)}(\theta^{(i+1)}) \right|_{\theta=\theta^{(i+1)}} = 0$$

with:

$$\left. \frac{\partial}{\partial \theta} V^{(i)}(\theta^{(i+1)}) \right|_{\theta=\theta^{(i+1)}} = \left. \frac{\partial}{\partial \theta} J_N(\theta) \right|_{\theta=\theta^{(i+1)}} + \left. \frac{\partial^2}{\partial \theta^2} J_N(\theta) \right|_{\theta=\theta^{(i+1)}} (\theta^{(i+1)} - \theta^{(i)})$$

and we obtain the condition for $\theta^{(i+1)}$: $\left. \frac{\partial}{\partial \theta} J_N(\theta) \right|_{\theta=\theta^{(i+1)}} + \left. \frac{\partial^2}{\partial \theta^2} J_N(\theta) \right|_{\theta=\theta^{(i+1)}} (\theta^{(i+1)} - \theta^{(i)}) = 0$

So the

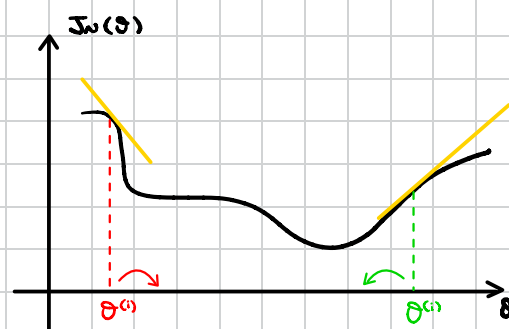
$$\theta^{(i+1)} = \theta^{(i)} - \left[\left. \frac{\partial^2}{\partial \theta^2} J_N(\theta) \right|_{\theta=\theta^{(i)}} \right]^{-1} \left. \frac{\partial}{\partial \theta} J_N(\theta) \right|_{\theta=\theta^{(i)}}$$

inverse of the Hessian matrix in $\theta^{(i)}$ gradient of $\theta^{(i)}$

Looking at the gradient, we have that:

- if it's positive $\Rightarrow \theta^{(i+1)} < \theta^{(i)}$
- if it's negative $\Rightarrow \theta^{(i+1)} > \theta^{(i)}$

We used an explicit form related to our signal



Computation of the gradient and the Hessian Matrix

Remembering that $J_N(\theta) = \frac{1}{N} \sum \epsilon(t|t-1, \theta)^2$ we derive it

gradient:

$$\frac{\partial}{\partial \theta} J_N(\theta) = \frac{d}{d\theta} \left[\frac{1}{N} \sum \epsilon(t|t-1, \theta)^2 \right] = \frac{1}{N} \sum \frac{d}{d\theta} \epsilon(t|t-1, \theta)^2 = \frac{2}{N} \sum \epsilon(t|t-1, \theta) \frac{\partial}{\partial \theta} \epsilon(t|t-1, \theta)$$

vector scalar vector

Hessian:

$$\begin{aligned} \frac{\partial^2}{\partial \theta^2} J_N(\theta) &= \frac{d}{d\theta} \left[\frac{\partial}{\partial \theta} J_N(\theta) \right] = \frac{d}{d\theta} \left[\frac{2}{N} \sum \epsilon(t|t-1, \theta) \frac{\partial}{\partial \theta} \epsilon(t|t-1, \theta) \right] \\ &= \frac{2}{N} \sum \frac{\partial}{\partial \theta} \epsilon(t|t-1, \theta) \left(\frac{\partial}{\partial \theta} \epsilon(t|t-1, \theta) \right)^T + \frac{2}{N} \sum \epsilon(t|t-1, \theta) \frac{\partial^2}{\partial \theta^2} \epsilon(t|t-1, \theta) \end{aligned}$$

square ≥ 0 (always) remove this because could be negative

We remove the second term to be sure to have a paraboloid with a minimum (and not a maximum).

Newton's Law (update rule for θ)

$$\theta^{(i+1)} = \theta^{(i)} - \left[\sum \left. \frac{\partial}{\partial \theta} \epsilon(t|t-1, \theta) \right|_{\theta^{(i)}} \left. \frac{\partial}{\partial \theta} \epsilon(t|t-1, \theta) \right|_{\theta^{(i)}}^T \right]^{-1} \sum \epsilon(t|t-1, \theta^{(i)}) \left. \frac{\partial}{\partial \theta} \epsilon(t|t-1, \theta) \right|_{\theta^{(i)}}$$

we used:

- $\epsilon(t|t-1, \theta^{(i)}) = y(t) - \hat{y}(t|t-1, \theta^{(i)})$ no need to further job, everything is in the data
- $\left. \frac{\partial}{\partial \theta} \epsilon(t|t-1, \theta) \right|_{\theta^{(i)}} = ?$

Computation of $\frac{\partial}{\partial \theta} \varepsilon(t|t-1, \theta)$

Being $\theta = [a_1 \dots a_m \ b_0 \dots b_{p-1} \ c_1 \dots c_n]^T$ the set of model parameters, we compute the prediction error writing it in the canonical form

$$\varepsilon(t|t-1, \theta) = \frac{A(z)}{C(z)} y(t) - \frac{B(z)}{C(z)} u(t-d) = \frac{1-a_1 z^{-1}-\dots-a_m z^{-m}}{1+c_1 z^{-1}+\dots+c_n z^{-n}} y(t) - \frac{b_0+b_1 z^{-1}+\dots+b_{p-1} z^{-(p-1)}}{1+c_1 z^{-1}+\dots+c_n z^{-n}} u(t-d)$$

So, when we derive $\varepsilon(t|t-1, \theta)$ by θ we obtain the vector

$$\frac{\partial}{\partial \theta} \varepsilon(t|t-1, \theta) = \left[\frac{\partial}{\partial a_1} \varepsilon(t|t-1, \theta) \dots \frac{\partial}{\partial b_0} \varepsilon(t|t-1, \theta) \dots \frac{\partial}{\partial c_n} \varepsilon(t|t-1, \theta) \dots \right]^T$$

That being $\theta = [a_1 \dots a_m \ b_0 \dots b_{p-1} \ c_1 \dots c_n]^T$ we derive for each element for θ : for example

$$\begin{aligned} \bullet \frac{\partial}{\partial a_1} \varepsilon(t|t-1, \theta) &= \frac{\partial}{\partial a_1} \left[\frac{1-a_1 z^{-1}-\dots-a_m z^{-m}}{C(z)} y(t) - \frac{B(z)}{C(z)} u(t-d) \right] && y(t) \rightarrow \boxed{\frac{-1}{C(z)}} \rightarrow \alpha(t) \\ &= \frac{\partial}{\partial a_1} \left[\frac{-a_1 z^{-1}}{C(z)} y(t) \right] = -\frac{z^{-1}}{C(z)} y(t) \triangleq \alpha(t-1) \\ \bullet \frac{\partial}{\partial a_2} \varepsilon(t|t-1, \theta) &= \frac{\partial}{\partial a_2} \left[\frac{-a_2 z^{-2}}{C(z)} y(t) \right] = -\frac{z^{-2}}{C(z)} y(t) \triangleq \alpha(t-2) \\ \bullet \frac{\partial}{\partial b_0} \varepsilon(t|t-1, \theta) &= \frac{\partial}{\partial b_0} \left[\frac{A(z)}{C(z)} y(t) - \frac{b_0+b_1 z^{-1}+\dots+b_{p-1} z^{-(p-1)}}{C(z)} u(t-d) \right] && u(t) \rightarrow \boxed{\frac{-1}{C(z)}} \rightarrow \beta(t-d) \\ &= \frac{\partial}{\partial b_0} \left[-\frac{b_0}{C(z)} u(t-d) \right] = -\frac{1}{C(z)} u(t-d) \triangleq \beta(t-d) \\ \bullet \frac{\partial}{\partial c_i} \varepsilon(t|t-1, \theta) &= \frac{\partial}{\partial c_i} \left[\frac{A(z)}{1+c_1 z^{-1}+\dots+c_n z^{-n}} y(t) - \frac{B(z)}{1+c_1 z^{-1}+\dots+c_n z^{-n}} u(t-d) \right] ? && \varepsilon(t|t-1, \theta) \rightarrow \boxed{\frac{-1}{C(z)}} \rightarrow \gamma(t) \end{aligned}$$

Computing $\frac{\partial}{\partial c_i} \varepsilon(t|t-1, \theta)$ is not straightforward since $C(z)$ is at the denominator. So we have to rewrite $\varepsilon(t|t-1, \theta)$ as

$$C(z) \varepsilon(t|t-1, \theta) = A(z) y(t) - B(z) u(t-d)$$

and now derive it by c_i

$$\begin{aligned} \frac{\partial}{\partial c_i} [C(z) \varepsilon(t|t-1, \theta)] &= \frac{\partial}{\partial c_i} [A(z) y(t) - B(z) u(t-d)] = 0 \\ &= \frac{\partial}{\partial c_i} C(z) \varepsilon(t|t-1, \theta) + C(z) \frac{\partial}{\partial c_i} \varepsilon(t|t-1, \theta) \end{aligned}$$

So the final derivation would be

$$\frac{\partial}{\partial c_i} \varepsilon(t|t-1, \theta) = -\frac{1}{C(z)} \frac{\partial}{\partial c_i} C(z) \cdot \varepsilon(t|t-1, \theta) = -\frac{1}{C(z)} z^{-i} \varepsilon(t|t-1, \theta) \triangleq \gamma(t-i)$$

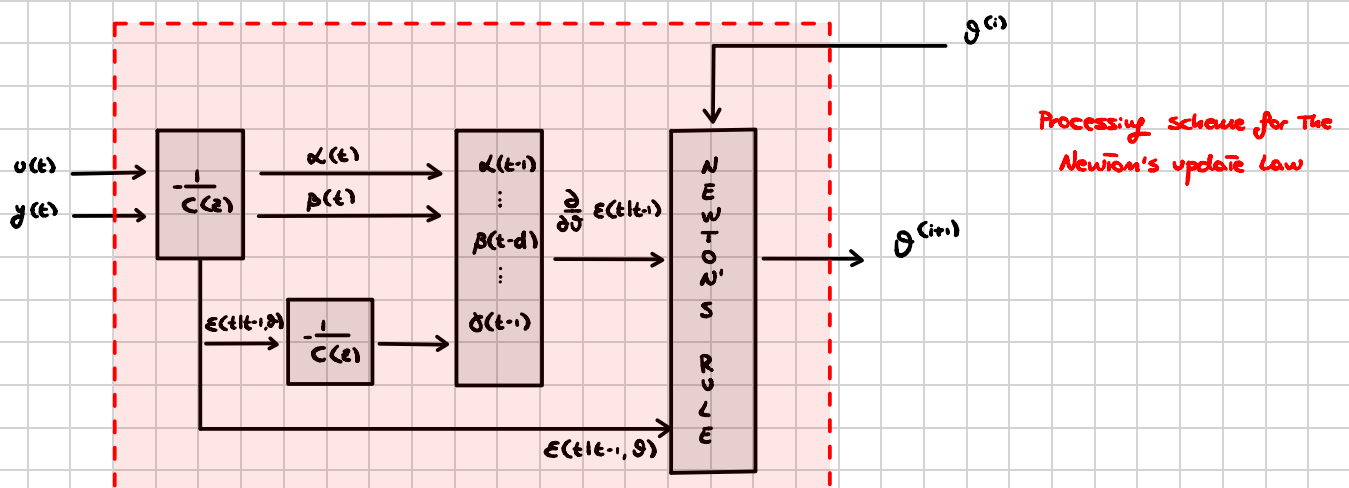
$\downarrow \frac{\partial}{\partial c_i} [1+c_1 z^{-1}+\dots+c_n z^{-n}]$

Final expression

$$\frac{\partial}{\partial \theta} \varepsilon(t|t-1, \theta) = [\alpha(t-1) \dots \alpha(t-m) \ \beta(t-d) \dots \beta(t-d-p+1) \ \gamma(t-1) \dots \gamma(t-n)]$$

where

$$\alpha = -1/C(z) \cdot y(t) \quad \beta = -1/C(z) \cdot u(t-d) \quad \gamma = -1/C(z) \cdot z^{-i} \varepsilon(t|t-1, \theta)$$



Remark: The Newton's rule is one of the possible way to solve this nonlinear optimization problem that takes advantage of the computation of the gradient and the second derivative of the cost.

- Newton's Rule $\theta^{(i+1)} = \theta^{(i)} - [\text{Hessian}]^{-1} \cdot \text{gradient}$
- gradient descent: $\theta^{(i+1)} = \theta^{(i)} - \eta \cdot \text{gradient}$ η scalar and fixed (step size)
- quasi-Newton's rule: $\theta^{(i+1)} = \theta^{(i)} - [\text{approximate Hessian}]^{-1} \cdot \text{gradient}$

→ Etc when we have removed the last piece from the Hessian: This is what we do! computationally

faster than Newton's, still more accurate but slower than GD