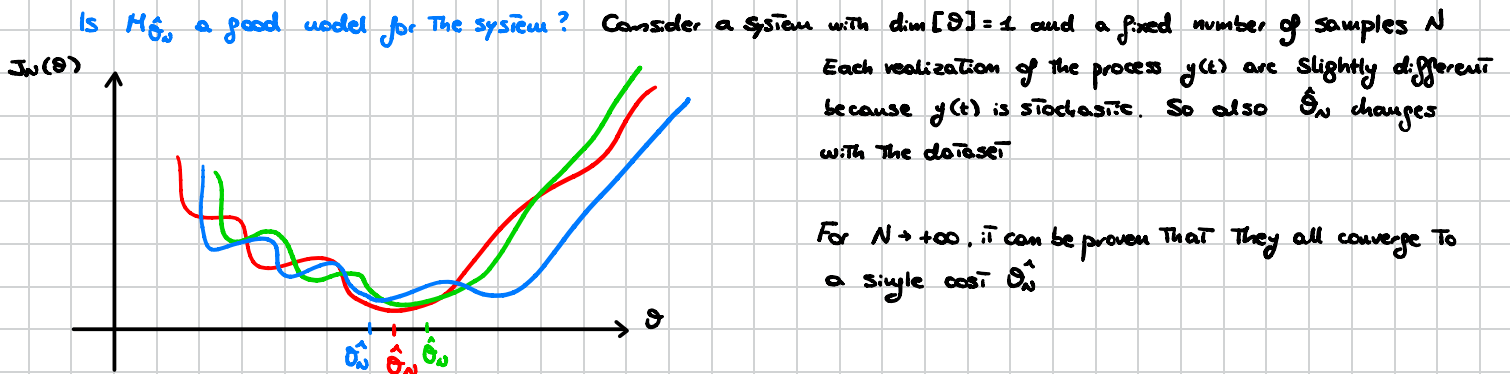


12/03 - Model Validation

Dataset: $D_N = \{(u(1), y(1)), \dots, (u(N), y(N))\}$
 Model class: $M_\Theta = \{M(\Theta), \Theta \in \Theta \subseteq \mathbb{R}^{n_\Theta}\}$
 M. param: $\hat{\Theta}_N = \arg\min_{\Theta} J_N(\Theta)$
 Avg squared error: $J_N(\Theta) = \frac{1}{N} \sum_{t=1}^N \epsilon(t|t-1, \Theta)^2$
 Prediction error: $\epsilon(t|t-1, \Theta) = y(t) - \hat{y}(t|t-1, \Theta)$



THEOREM: Under current assumption, as $N \rightarrow +\infty$, the predicted cost $J_N(\Theta, s) \rightarrow \bar{J}(\Theta) = E_s[E(t|t-1, \Theta, s)^2]$. Moreover, by letting $\Delta = \{\Theta^*, \bar{J}(\Theta^*) \leq \bar{J}(\Theta), \forall \Theta\}$ be the set of global minimum points of $\bar{J}(\Theta)$, we have

$$\hat{\Theta}_N(s) \rightarrow \Delta \quad \text{for } N \rightarrow +\infty$$

with probability 1.

Corollary: if $\Delta = \{\Theta^*\}$, i.e. $\bar{J}(\Theta)$ has one unique minimum point, $\hat{\Theta}_N(s) \rightarrow \Theta^*$ wp.1 for $N \rightarrow +\infty$.

Are we happy with $\hat{\Theta}_N \rightarrow \Theta^*$? What we are really interested in is to guarantee that if $S \in M_\Theta$ ($\exists \Theta^0: S \equiv M(\Theta^0)$), then $\hat{\Theta}_N \rightarrow \Theta^0$.

Theorem: $\Theta^0 \in \Delta$ (so, if $\Delta = \{\Theta^*\}$ Singleton, so $\Theta^0 = \Theta^*$)

Proof:

$$\begin{aligned}
 \text{Prediction error: } \epsilon(t|t-1, \Theta) &= y(t) - \hat{y}(t|t-1, \Theta) \\
 &= y(t) - \hat{y}(t|t-1, \Theta^0) + \hat{y}(t|t-1, \Theta^0) - \hat{y}(t|t-1, \Theta) \\
 &\quad \text{e(t)} \quad \text{optimal predictor}
 \end{aligned}$$

$y(t) = ay(t-1) + e(t) \Rightarrow \hat{y}(t|t-1) = ay(t-1) \Rightarrow e \cdot y \cdot \hat{y} = e(t) \Rightarrow \text{VAR}[e(t-1)] = \text{VAR}[e(t)]$
 for the 1-step ahead predictor

In this way the identification criteria $\bar{J}(\Theta)$, so the variance of the prediction error $\epsilon(t|t-1, \Theta)$ will be

$$\begin{aligned}
 \bar{J}(\Theta) &= \text{VAR}[E(t|t-1, \Theta)] = E[\epsilon^2(t|t-1, \Theta)] \\
 &= E[(e(t) + \hat{y}(t|t-1, \Theta^0) - \hat{y}(t|t-1, \Theta))^2] \\
 &= E[e(t)^2] + E[(\hat{y}(t|t-1, \Theta^0) - \hat{y}(t|t-1, \Theta))^2] + \\
 &\quad \underline{2E[e(t)(\hat{y}(t|t-1, \Theta^0) - \hat{y}(t|t-1, \Theta))]} = 0
 \end{aligned}$$

So the asymptotic cost

$$\bar{J}(\Theta) = \lambda^2 + E[(\hat{y}(t|t-1, \Theta^0) - \hat{y}(t|t-1, \Theta))^2]$$

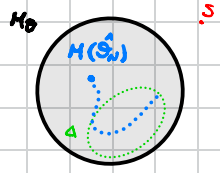
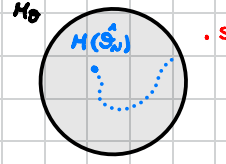
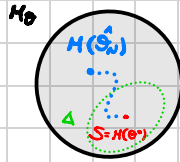
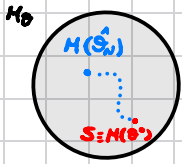
where $E[(\hat{y}(t|t-1, \Theta^0) - \hat{y}(t|t-1, \Theta))^2]$ is function of $\Theta \geq 0$ so the minimum will be found for $\Theta = \Theta^0$.

So $\bar{J}(\Theta^0) \leq \bar{J}(\Theta) \quad \forall \Theta$ which means that $\Theta^0 \in \Delta$.

predictors computed based on data until $t-1$

4 Possible scenarios:

We can encounter 4 possible scenarios:



- a) $S \in H_0$ and Δ is singleton ($S \equiv H(\theta^*)$, $\Delta = \{\theta^*\}$)
- b) $S \in H_0$ but Δ is not singleton (Δ contains more than 1 global optimum). Not possible if everything is well posed, but it might happen
- c) $S \notin H_0$ and Δ is singleton
- d) $S \notin H_0$ and Δ is not singleton

Selecting H_0 is the most critical part of the process

Model order selection

Selecting $H_0 = \{H(\theta) : \theta \in \Theta \subseteq \mathbb{R}^{n_\theta}\}$. We cannot guarantee that $S \in H_0$ because S is unknown and we don't know how to enlarge H_0 in case $S \notin H_0$. For simplicity, we will consider $m = p$ (only 1 degree of freedom and simplify the process). How to select m properly (scalar optimization problem)?

Let us take the following procedure

set $m = 1$

repeat until $m = m_{\max}$

$$H_0^{(m)} = \{H(\theta) : \theta \in \Theta^{(m)} \subseteq \mathbb{R}^{m_\theta}\}$$

$$\hat{\theta}_N^{(m)} = \arg \min_{\theta} J_N^{(m)}(\theta)$$

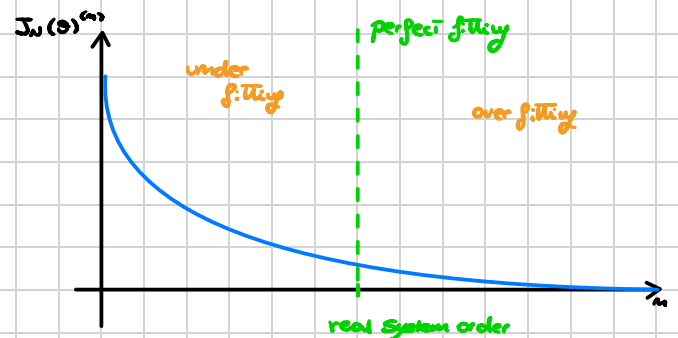
$m := m + 1$

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

Consider by starting with: $A(z) = 1 - a_1 z^{-1}$ $B(z) = b_0 + b_1 z^{-1}$ $C(z) = 1 + c_1 z^{-1}$ and then add terms to each component.

This is not suitable for the selection of m , since

This point cannot be detected just by looking at the behaviour of $J_N(\hat{\theta}_N^{(m)})$ as function of m



3 different criteria for model order selection

- whiteness Test on residual
- cross validation
- identification with model order penalties

Whiteness Test on Residuals (Anderson's Test)



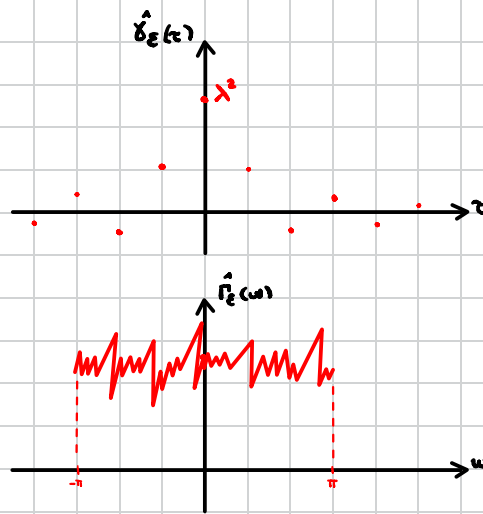
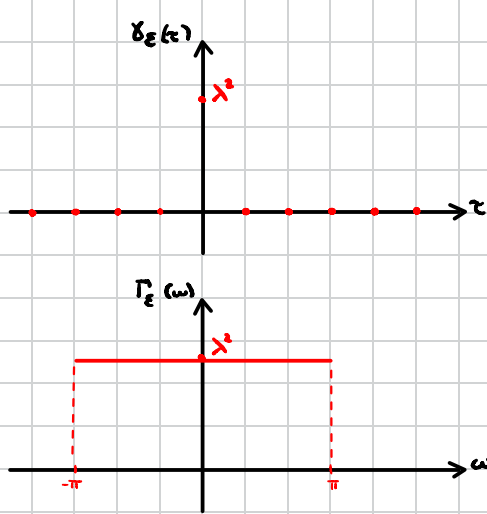
We know that, if $S \in H_0$, for a certain model order m , $\hat{\theta}_N^{(m)} \rightarrow \theta^*$, and

$$\epsilon(t|t-1, \hat{\theta}_N^{(m)}) \rightarrow e(t) \sim WN(0, \lambda^2)$$

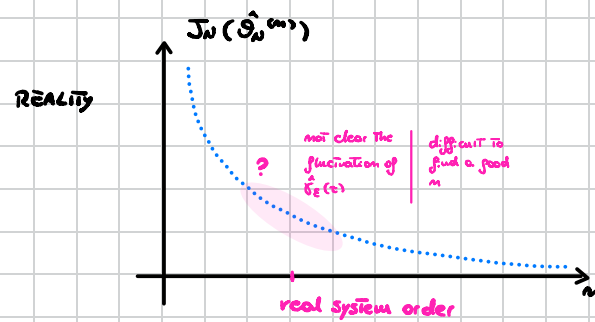
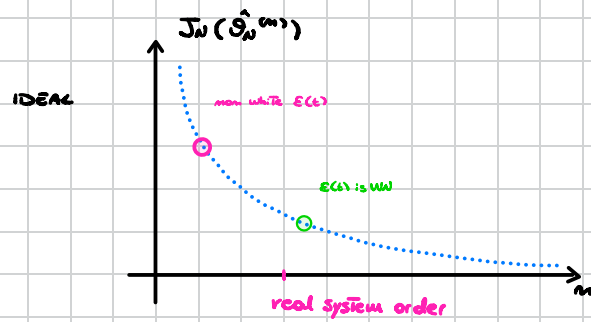
Then, for large N :

$$\epsilon(t|t-1, \hat{\theta}_N^{(m)}) = y(t) - \hat{y}(t|t-1, \hat{\theta}_N^{(m)})$$

should be a realization of a WN



If $E(t)$ is WW, we expect the real values. However, we obtain some fluctuation that make difficult to recognize a WW.

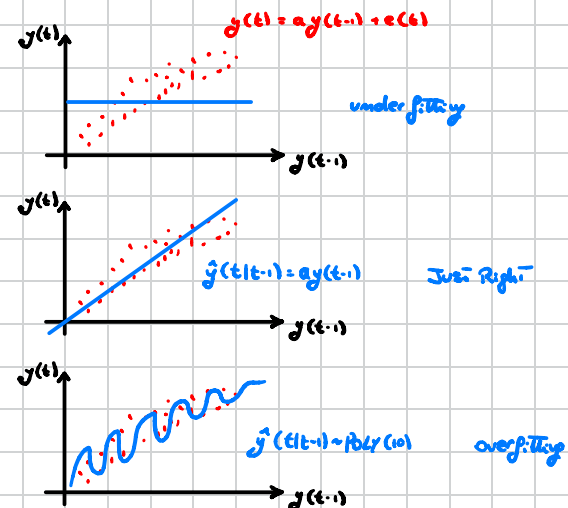
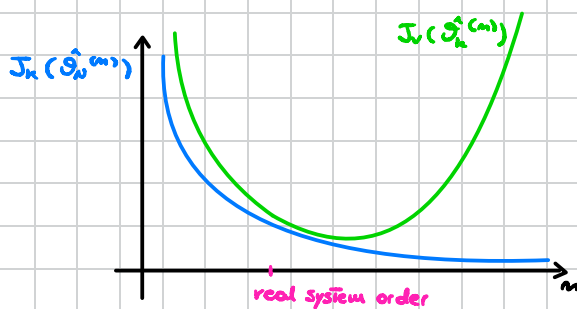


Cross Validation

Consider the dataset $D_N = \{(u(1), y(1)), \dots, (u(m), y(m))\}$ and the cost function we want to minimize $J_N(\theta) = \frac{1}{N} \sum (y(t) - \hat{y}(t|t-1, \theta))^2$. Just looking at the dataset we can subdivide it into training set and validation set. We can follow the procedure:



For $n=1$ To $m=m_{max}$
 find $\hat{\theta}_n^{(m)} = \arg \min_{\theta} J_n^{(m)}(\theta)$
 evaluate $J_v(\hat{\theta}_n^{(m)}) = \frac{1}{N-k} \sum_{t=k+1}^N (y(t) - \hat{y}(t|t-1, \hat{\theta}_n^{(m)}))^2$
 Choose m (optimal model order) as the one minimizing J_v
 $\hat{m} = \arg \min_m J_v(\hat{\theta}_n^{(m)})$



In ML, **k-fold cross validation** is often used: divide the dataset in k blocks, using $1/k$ of it as validation and the other as training. This is not doable here since the system is dynamical and we cannot attach different time windows (this would lead to "fake" temporal relations). Same problem for validation.

Identification costs with model order penalties

Using Cross Validation we have split the dataset and obtained the same cost for Training and validation

Adv: different costs for model order selection and parameter selection $\rightarrow$ Single DS

Instead of minimizing $J_N(\theta) = \frac{1}{N} \sum (y(t) - \hat{y}(t|t-1, \theta))^2$ we can minimize 3 alternatives

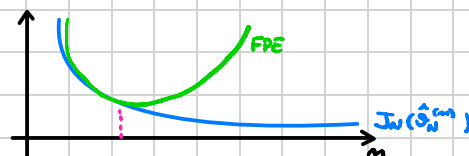
- 1) Final prediction error (FPE)
- 2) Akaike information criterion (AIC)
- 3) Minimum description length (MDL)

FPE: our ideal objective is $\min_{\theta} E[J_N(\theta)]$ but is not affordable. So, FPE is a numerical approximation of this

$$FPE(m) = \frac{N+m}{N-m} J_N(\hat{\theta}_N^{(m)})$$

AIC: $AIC(m) = \underbrace{\log(J_N(\hat{\theta}_N^{(m)}))}_{\text{decrease with } m} + 2 \underbrace{\frac{m}{N}}_{\text{penalty}}$

MDL: $MDL(m) = \underbrace{\log(J_N(\hat{\theta}_N^{(m)}))}_{\text{decrease with } m} + \underbrace{\log(N) \frac{m}{N}}_{\text{penalty}}$



FPE vs AIC?

$$\log(FPE) = \log\left(\frac{N+m}{N-m} J_N(\hat{\theta}_N^{(m)})\right) = \log\left(\frac{1+\frac{m}{N}}{1-\frac{m}{N}} J_N(\hat{\theta}_N^{(m)})\right)$$

$\frac{m}{N}$ is small

$$\Rightarrow \log\left(\frac{1+x}{1-x}\right) \approx 2x \text{ as } x \rightarrow 0$$

$$= \log(1 + \frac{m}{N}) - \log(1 - \frac{m}{N}) + \log(J_N(\hat{\theta}_N^{(m)}))$$

$$\approx \frac{m}{N} - (-\frac{m}{N}) + \log(J_N(\hat{\theta}_N^{(m)}))$$

$$= \log(J_N(\hat{\theta}_N^{(m)})) + 2 \frac{m}{N} = AIC$$

$\Rightarrow$ So AIC and FPE return the same optimal $\hat{m}$

AIC vs MDL? Being

$$AIC = \log(J_N(\hat{\theta}_N^{(m)})) + 2 \frac{m}{N}$$

$$MDL = \log(J_N(\hat{\theta}_N^{(m)})) + \log(N) \frac{m}{N}$$

if $\log(N) \geq 2 \Rightarrow$ The model order is penalized more by MDL $\Rightarrow N > 8$ (basically always since $N \sim 10000$)

If $S \in M_0$ and M_0 is the set of ARX models, MDL is right. In the general case, $S \in M_0$ we usually prefer to slightly overfitting and we use AIC