

THE VECTOR SPACE OF RANDOM VARIABLES.

Now, the set of real random variables defined on a probability space with the field of real numbers constitute a vector space with $X + Y$ the vector addition of two random variables X, Y and aX the multiplication a vector (r.v X) by a scalar a (easy to check). However, the vector spaces of r.v.s are somewhat more complicated than R^n or C^n . Recall that R^n and C^n are all finite dimensional spaces so that it is possible to specify every vector as a linear combination of a number of basis vectors. In the case, of r.v.s, unless the sample space is finite the vector space becomes infinite dimensional.

INNER PRODUCT

An additional property that a vector space over a scalar field may have is the existence of an inner product. This is a function mapping vector pairs to the scalar fields.

An inner product which is denoted as $\langle \cdot, \cdot \rangle$ must satisfy the following properties:

1. Existence: $\langle X, Y \rangle$ is defined for all vectors X, Y
2. Linearity: $\langle X, \alpha Y + \beta Z \rangle = \alpha \langle X, Y \rangle + \beta \langle X, Z \rangle$
3. Hermitian symmetry: $\langle X, Y \rangle = \langle Y, X \rangle^H$
4. Positive definiteness: $\langle X, X \rangle \geq 0$ if $X \neq 0$

* A vector space with an inner product is called inner product space

* R^n with the well known dot product is an inner product space

* For the vector space of random variables a valid inner product is

$$\langle X, Y \rangle \triangleq E\{X^*Y\} = \text{Cov}(X, Y) \text{ (for zero mean r.v.s)}$$

as it can be easily verified (assuming of course that $E\{|X|^2\} < \infty$ for all r.v.s X in the space).

NORM or DISTANCE

Recall that the notion of length or distance or norm of a vector is define as

$$||X|| = \langle X, X \rangle^{1/2}$$

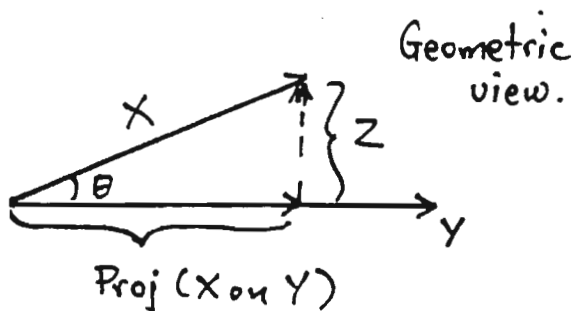
and satisfies the following properties:

1. $||aX|| = |a| ||X||$
2. $||X|| > 0$ unless $X = 0$ where $||0|| = 0$
3. $|\langle X, Y \rangle| \leq ||X|| ||Y||$ (Cauchy-Schwarz inequality)
4. $||X + Y|| \leq ||X|| + ||Y||$ (Triangle inequality)

PROJECTION

Given two vectors X and Y , the projection of X on Y is given by:

$$\text{Proj}(X \text{ on } Y) = \langle X, \frac{Y}{||Y||} \rangle \frac{Y}{||Y||}$$



For vector spaces over $\mathbb{R}$ only we can write (by inspection from the geometric view)

$$||X|| \cos \theta = \langle X, \frac{Y}{||Y||} \rangle = \langle X, Y \rangle \frac{1}{||Y||} \Rightarrow \langle X, Y \rangle = ||X|| ||Y|| \cos \theta$$

(Note that $\text{proj}(X, Y)$ is nothing else but the vector in the form of Y which is closest to X in the norm sense)

More general in the vector space of random variables (assume zero mean r.v for the moment)

$$\langle X, Y \rangle = E\{X^*Y\} = \text{cov}(X, Y)$$

$$\|X\|^2 = E\{X^*X\} = \text{var}(X)$$

And from the familiar relation $\rho = \frac{\text{cov}(X, Y)}{\sqrt{\text{var}(X)\text{var}(Y)}}$

$$\Rightarrow \boxed{\langle X, Y \rangle = \rho \cdot \|X\| \|Y\|} \Rightarrow \langle X, \frac{Y}{\|Y\|} \rangle = \rho \cdot \|X\|$$

(i.e., $\cos\theta$ is the correlation coefficient in vector spaces over R)

ORTHOGONALITY

Two vectors X and Y are called orthogonal iff $\langle XY \rangle = 0$. In vector spaces of random variables this is equivalent to $\rho = 0$. Looking back into the $\text{proj}(X \text{ on } Y)$ we see that (zero mean r.v.s)

$$\begin{aligned} E\{ZY^*\} &= E\{(X - \text{proj}(X \text{ on } Y))Y^*\} = E[(X - \langle X, \frac{Y}{\|Y\|} \rangle \frac{Y}{\|Y\|})Y^*] \\ &= E\{(X - \rho \frac{\|X\|}{\|Y\|} Y)Y^*\} = E\{XY^*\} - \rho \frac{\|X\|}{\|Y\|} E\{YY^*\} \\ &= \rho\sigma_X\sigma_Y - \rho \frac{\sigma_X}{\sigma_Y} \sigma_Y^2 = 0 \Rightarrow \\ E\{ZY^*\} &= 0 \Rightarrow Z, Y \text{ orthogonal} \end{aligned}$$

MEAN SQUARE ESTIMATION

Let us consider the problem where we are given a random variable Y , and we are asked to find the best approximation to Y using a particular model. Obviously, the solution to the problem will depend on the criterion of optimality chosen. We would like to consider this problem using the criterion of minimizing the mean square error of the difference between Y and its approximation. (Mean square estimation of Minimum mean square error model (MSE)). The results are the basis for a great number of applications in signal processing.

A) Consider the MSE problem where we want to approximate (estimate) a random variable Y by a constant c :

Choose c so that to minimize $\|Y - c\|^2$

$$\text{Now, } \|Y - c\|^2 = E\{|Y - c|^2\} = E\{|Y|^2\} + |c|^2 - 2\text{Re}[c^* E\{Y\}]$$

and by placing

$$\frac{d}{dc} \|Y - c\|^2 = 0 \Rightarrow c^* - [E\{Y\}]^* = 0 \Rightarrow c = E\{Y\}$$

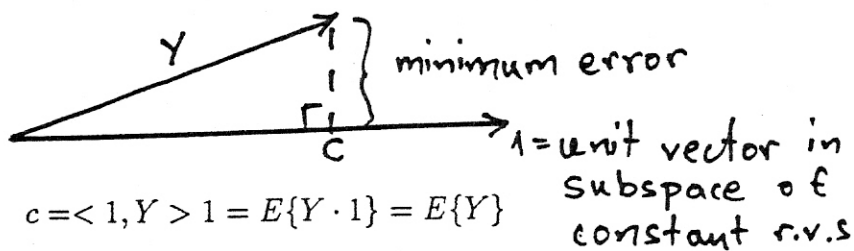
In other words, the best constant approximation of a random variable is its mean

$$\text{Also, } \min \|Y - c\|^2 = E\{|Y - E\{Y\}|^2\} = \text{Var}\{Y\} = \sigma_y^2$$

Geometric interpretation

Let consider a constant as a r.v. which maps the entire sample space to a constant value. Then the entire collection of such r.v.s constitute a 1-D subspace of the vector space of r.v.s

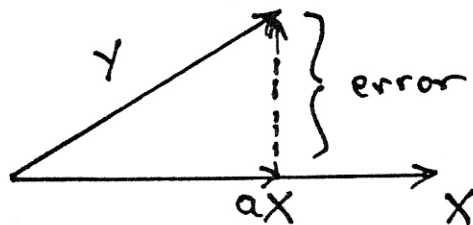
Then geometrically the best estimate is the projection of Y on the subspace



Note that the error $Y - c$ such that

$$\langle Y - c, 1 \rangle = 0$$

B) Consider now the problem where we want to approximate Y with a scaled version of another r.v. X i.e., aX . Then



minimizing $\|Y - aX\|^2 \Rightarrow \text{error orthogonal to } aX$

$$\Rightarrow \langle Y - aX, aX \rangle = 0 \Rightarrow E\{(Y - aX)aX^*\} = 0$$

$$\Rightarrow aE\{X^*Y\} - a^2E\{|X|^2\} = 0 \Rightarrow E\{X^*Y\} = a \cdot E\{|X|^2\}$$

$$\Rightarrow a = \frac{E\{X^*Y\}}{E\{|X|^2\}} \xrightarrow{\text{zero mean process}} a = \frac{\text{Cov}(X, Y)}{\text{Var}(X)} = \rho \frac{\sigma_Y}{\sigma_X}$$

and

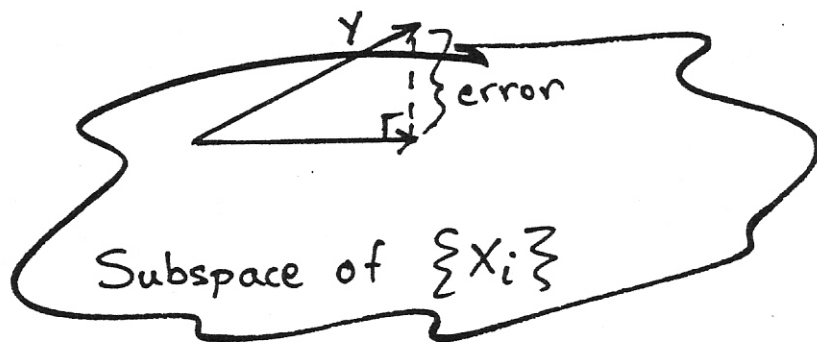
$$\min \|Y - aX\|^2 = \sigma_Y^2 + \rho^2 \sigma_Y^2 - \sigma_X^2 2a \text{Re}(a)^{\text{real } a} \sigma_Y^2 (1 - \rho^2)$$

C) The most general form of this type of problem is to approximate the r.v. Y by a linear combination of n r.v.s $X_1, X_2, \dots, X_n$, i.e., to choose $a_1, a_2, \dots, a_n$ so that.

$$\|Y - (a_1X_1 + \dots + a_nX_n)\|^2$$

is minimized (Linear mean square estimation problem (LMSE)).

To solve this problem we note that the set of r.v. of the form $a_1X_1 + \dots + a_nX_n$ constitute a subspace of the vector space of r.v.s. Geometrically, the solution on of minimizing the magnitude squared of the error quantity $(Y - \sum_{i=1}^n a_iX_i)$ is obtained when the error becomes orthogonal to the approximation $\sum_{i=1}^n a_iX_i$. i.e.,



Thus, $\min \|Y - (a_1X_1 + \dots + a_nX_n)\|$ is equivalent to

$$\langle Y - \sum_{i=1}^n a_iX_i, \sum_{i=1}^n a_iX_i \rangle = 0$$

This is the projection theorem for the LMSE problem

$$\Rightarrow E\{(Y - \sum_{i=1}^n a_iX_i) \sum_{j=1}^n a_j^* X_j^*\} = 0 \Rightarrow \sum_{j=1}^n a_j^* [E\{X_j^* \cdot (Y - \sum_{i=1}^n a_iX_i)\}] = 0$$

Since this must be true for all $\{a_j\}_{j=1, \dots, n}$

$$\Rightarrow E\{X_j^* \cdot (Y - \sum_{i=1}^n a_iX_i)\} = 0 \Rightarrow E\{X_j^* Y\} = \sum_{i=1}^n a_i E\{X_j^* X_i\}, \quad j = 1, 2, \dots, n$$

For X,Y jointly WSS r.v.s

$$\boxed{\sum_{i=1}^n a_i R_x(i-j) = R_{xy}(-j) \quad j = 1, \dots, n} \quad \text{Wiener - Hopf equations}$$

or

$$\begin{bmatrix} R_x(0) & R_x(1) & \dots & R_x(n-1) \\ R_x(-1) & R_x(0) & \dots & . \\ R_x(-2) & R_x(-1) & \dots & . \\ \vdots & \vdots & \vdots & \vdots \\ R_x(-n+1) & \dots & \dots & R_x(-1)R_x(0) \end{bmatrix} \begin{bmatrix} a_1 \\ a_2 \\ \vdots \\ a_n \end{bmatrix} = \begin{bmatrix} R_{xy}(-1) \\ R_{xy}(-2) \\ \vdots \\ R_{xy}(-n) \end{bmatrix} \quad \text{Normal equations}$$

or

$$\mathbf{R}_x \cdot \mathbf{a} = \mathbf{r}_{xy}$$

Thus,

$$\boxed{\mathbf{a} = \mathbf{R}_x^{-1} \mathbf{r}_{xy}}$$

It is easy to show that $\text{Min} \|Y - \sum_{i=1}^n a_i X_i\| = \sigma_y^2 - \mathbf{r}_{xy}^T \mathbf{R}_x^{-1} \mathbf{r}_{xy} = \sigma_y^2 - \rho_{xy}^2$

D) Let us consider now the more general problem where instead of a linear combination of r.v.s, we employ a more general model. For example we might seek a function $g(X)$ so that $\|Y - g(X)\|^2$ is minimized. Recall that

$$E\{|Y - g(X)|^2\} = E\{E\{|Y - g(X)|^2/X\}$$

Therefore, the problem becomes equivalent to minimizing $E\{|Y - g(X)|^2/X = x\}$. Then, from the results of case A, we conclude that the absolute mean square error is minimized when

$$\boxed{g(X) = E\{Y/X\}}$$

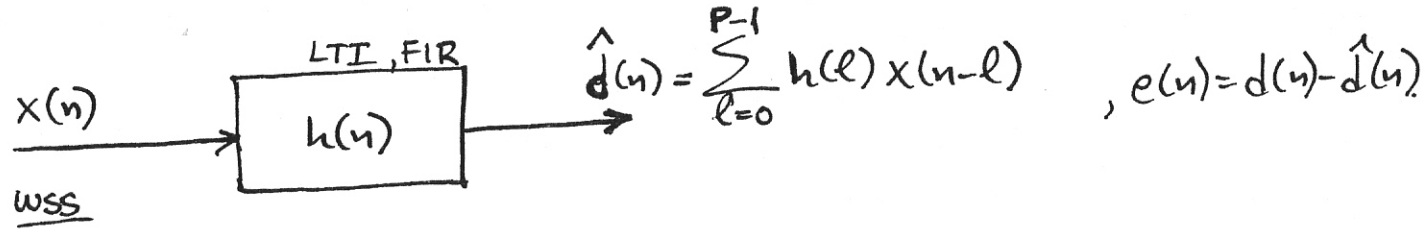
WIENER FILTERING PROBLEMS

We observe a possibly non-stationary process $X(n)$ over an interval I . We want to find the LMMSE of a process $d(n)$ (desired response) based on a linear combination of the r.v.s that constitute $X(n)$. The process $X(n)$ is related to some other "signal" process $S(n)$

PROBLEM

- 1) Filtering of signal in noise : $X(n) = S(n) + n(n)$, $d(n) = S(n)$
- 2) Prediction of signal in noise : $X(n) = S(n) + n(n)$, $d(n) = S(n+p)$, $p > 0$
- 3) Smoothing of signal in noise : $X(n) = S(n) + n(n)$, $d(n) = S(n-q)$, $q > 0$
- 4) Linear prediction : $X(n) = S(n-1)$, $d(n) = S(n)$
- 5) General nonlinear problem : $X(n) = G(S(n), n(n))$, $d(n) = S(n)$
- 6) Deconvolution : $X(n) = S(n) * h(n) + n(n)$, $d(n) = S(n)$

WIENER-HOPF EQUATION - STATIONARY SIGNALS



Then, minimization of $E\{|d(n) - \hat{d}(n)|^2\}$ results in

$$\left[\sum_{l=0}^{P-1} h(l) \cdot R_x(i-l) = R_{xd}(-i) \right] , \quad i = 0, 1, \dots, P-1.$$

where $R_x(m) \triangleq E\{x^*(n) x(n+m)\}$, $R_{xd}(m) \triangleq E\{x^*(n) d(n+m)\}$

Minimum mean square error: $\sigma_e^2 = \sigma_d^2 - \sum_{l=0}^{P-1} h(l) \cdot R_{xd}(l) = R_d(0) - \sum_{l=0}^{P-1} h(l) R_{xd}(l)$

Solution:

$$\begin{bmatrix} R_x(0) & R_x(-1) & \dots & R_x(1-P) \\ R_x(1) & R_x(0) & & \\ \vdots & & \ddots & \\ R_x(P-1) & \dots & \dots & R_x(0) \end{bmatrix} \begin{bmatrix} h(0) \\ h(1) \\ \vdots \\ h(P-1) \end{bmatrix} = \begin{bmatrix} R_{xd}(0) \\ R_{xd}(1) \\ \vdots \\ R_{xd}(P-1) \end{bmatrix}$$



WIENER-HOPF EQUATION - NON-CAUSAL CASE

$$\hat{d}(n) = \sum_{l=-\infty}^{\infty} h(l) x(n-l) \Rightarrow \sum_{l=-\infty}^{\infty} h(l) R_x(i-l) = R_{xd}(i) \Rightarrow$$

$$\Rightarrow \sum_{l=-\infty}^{\infty} h(l) * R_x(l) = R_{xd}(i) \Rightarrow h(l) * R_x(l) = R_{xd}(l)$$

$$\Rightarrow \boxed{H(z) = \frac{S_{dx}(z)}{S_x(z)}}$$

$$\left\{ \begin{array}{l} \text{based on definition} \\ R_{xd}(m) = E\{x^*(n) d(n+m)\} \end{array} \right\}$$

The non-causal equation provides insight into the operation of the Wiener filter.

Causal case: Assuming that $S_x(z)$ can be written as $S_x(z) = K H_{ca}(z) H_{ca}^*(1/z^*)$ where $H_{ca}(z)$ is a causal stable system with causal stable inverse.
 (h(l)=0, l<0)

Then, $H(z) = \frac{1}{K H_{ca}(z)} \left[\frac{S_{dx}(z)}{H_{ca}^*(1/z^*)} \right]_+$ where $[]_+$ corresponds only to the causal part of the response.



SPECTRAL FACTORIZATION

Given a random process $x(n)$ with continuous PSD $S_x(\omega)$:

If $\left[\int_{-\pi}^{\pi} |H_y S_x(\omega)| d\omega < \infty \right]$ then we can write. $\left[S_x(\cdot) = K \cdot H_{ca}(z) H_{ca}^*(1/z^*) \right]$

where, K is a positive constant, $H_{ca}(z)$ is a causal stable system with causal stable inverse.

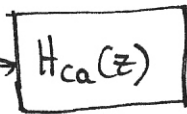
✿ If $S_x(z)$ is a rational polynomial, $H_{ca}(z)$ is minimum phase, $H_{ca}^*(1/z^*)$ maximum phase systems.

✿ The above condition is known as the Paley-Wiener condition.

Modeling of $x(n)$:

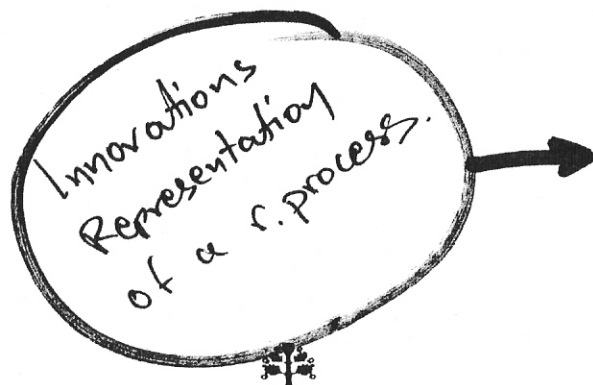
"innovations process"
 $w(n)$

white noise
 $S_w(z) = K$

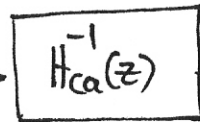


$x(n)$: "regular process"

$$S_x(z) = S_w(z) \cdot H_{ca}(z) \cdot H_{ca}^*(1/z^*)$$

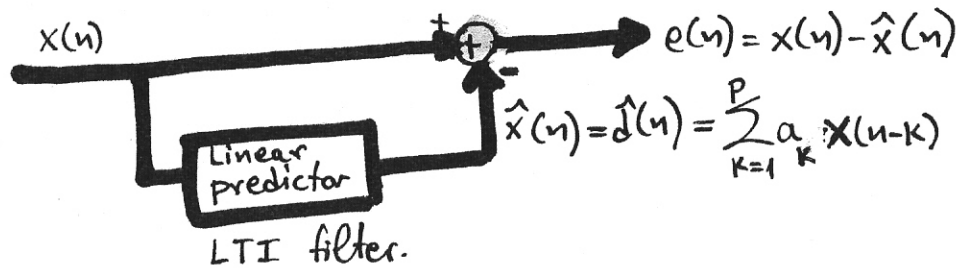


$x(n)$



$w(n)$: innovations process
(white noise)

LINEAR PREDICTION (FORWARD)



Objective: predict value of $x(n)$ at time n based on P past observations at $n-1, n-2, \dots, n-P$

$e(n)$: forward prediction error. $= x(n) - \sum_{k=1}^P a_k x(n-k)$

$d(n) = x(n)$

Minimization of $E\{|d(n) - \hat{d}(n)|^2\}$

results in

① $\left[\sum_{k=1}^P a_k R_x(i-k) = -R_x(i) \right], i=1, 2, \dots, P$

② $\sigma_e^2 = R_x(0) - \sum_{k=1}^P a_k R_x(k) = R_x(0) - \sum_{k=1}^P a_k R_x(k)$



①+② $\rightarrow$ "Normal Equations"

Whitening property of prediction filter.

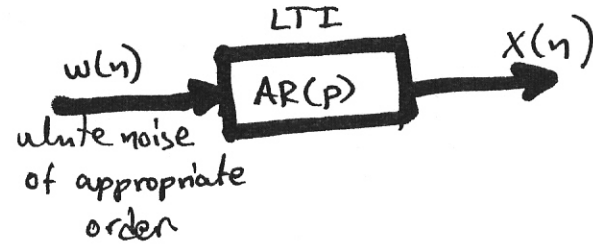
Calculation of $\{a_k\}$ is based on the Correlation properties of $x(n)$. The prediction filter estimates $x(n)$ by utilizing all correlation information between $x(n)$ and past observations. Thus, the minimum mean square error with minimum σ_e^2 corresponds to a white sequence. $e(n)$

CAUSAL AR modeling of a r. process

Given a r. pr. $x(n)$ generated from an AR(p) stable model.

↑
(autoregressive)

$$\left[x(n) + \sum_{k=1}^P a_k x(n-k) = w(n) \right]$$



* Multiply both sides by $x^*(n-l)$ and take expectation. Then

$$E\{x(n)x^*(n-l)\} + \sum_{k=1}^P a_k E\{x(n-k)x^*(n-l)\} = E\{w(n)x^*(n-l)\} \Rightarrow R_x(l) + \sum_{k=1}^P a_k R_x(l-k) = \begin{cases} \sigma_w^2, & l=0 \\ 0, & l>0 \end{cases}$$

This is nothing else but the normal equations

