

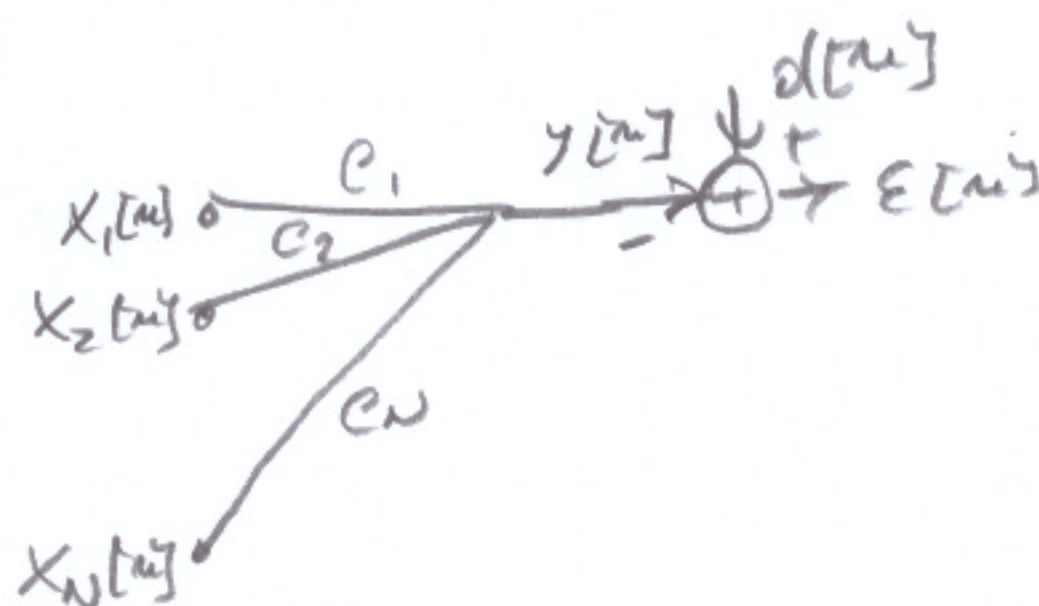
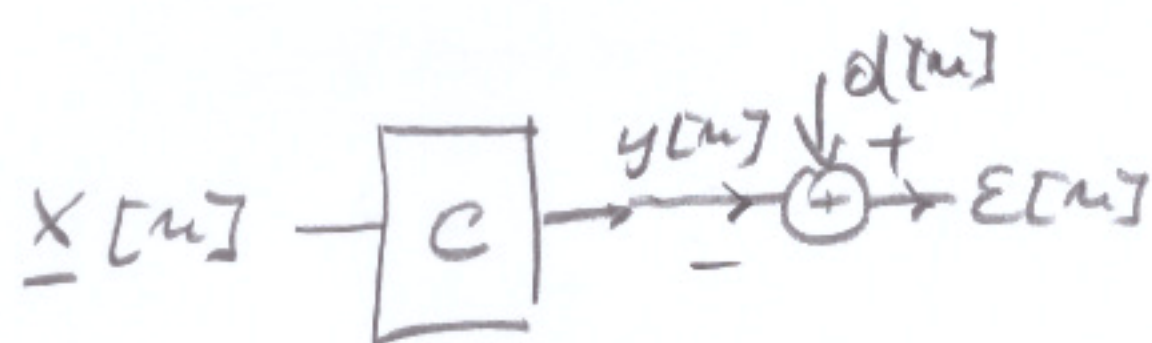
FILTRI A MINIMO
ERRORE QUADRATICO
(cont.)

Lezione di Telecomunicazioni.
Prof. FRANCESCO A.N. PALMIGRI

CORSO DI "SIGNAL PROCESSING AND DATA FUSION"
AA. 2022-23

CASO VETTORIALE GENERICO

Il paradigma non cambia se il vettore $\underline{x}[n]$ proviene da un processo o segnale vettoriale (e quindi non necessariamente una finestra che scorre su un segnale)



La soluzione è ovviamente

$$\begin{cases} \underline{c}^0 = \underline{R}_x^{-1} \underline{\zeta}_{dx} \\ \mathcal{E}(\underline{c}^0) = \mathcal{E}_d - \underline{\zeta}_{dx}^T \underline{R}_x^{-1} \underline{\zeta}_{dx} \end{cases} \quad \begin{cases} \underline{R}_x = E[\underline{x}[n] \underline{x}^T[n]] \\ \underline{\zeta}_{dx} = E[d[n] \underline{x}[n]] \\ \mathcal{E}_d = E[d^2[n]] \end{cases}$$

TRAINING SET

$$\{(d[n], \underline{x}[n]), n=1, \dots, L\}$$

$\Downarrow$
 ottieni $\underline{R}_x$ e $\underline{\zeta}_{dx}$

$\Downarrow$

Risolvi $\underline{c}^0 = \underline{R}_x^{-1} \underline{\zeta}_{dx}$

$$(*) \quad \mathcal{E}(\underline{c}^0) = \mathcal{E}_d - \underline{\zeta}_{dx}^T \underline{R}_x^{-1} \underline{\zeta}_{dx}$$

$$\begin{cases} \underline{R}_x = \frac{1}{L} \sum_{n=1}^L \underline{x}[n] \underline{x}^T[n] \\ \underline{\zeta}_{dx} = \frac{1}{L} \sum_{n=1}^L \underline{x}[n] d[n] \\ \mathcal{E}_d = \frac{1}{L} \sum_{n=1}^L d^2[n] \end{cases}$$

le prestazioni nel training set sono quelle ottenute stimando $\underline{R}_x$, $\underline{\zeta}_{dx}$ e $\mathcal{E}_d$ dal training set.

VALIDATION SET

le prestazioni nel validation set sono ancora date dalla formula (*) ma con $\underline{R}_x$, $\underline{\zeta}_{dx}$, $\mathcal{E}_d$ stimate dal validation set.

CORRELATIONS OR COVARIANCES AND THE ADDITION OF A BIAS

In solving our least squares problems
in considering the cost

$$E[(d-y)^2]$$

we have ignored an important aspect
of the problem. It may happen that
 d and/or x have non-zero means.
therefore $y = \underline{c}^T \underline{x}$ will have a non-zero
mean and the defined cost may be
strongly dominated by a search for
the filter to match the mean of d ("mean
choosing")

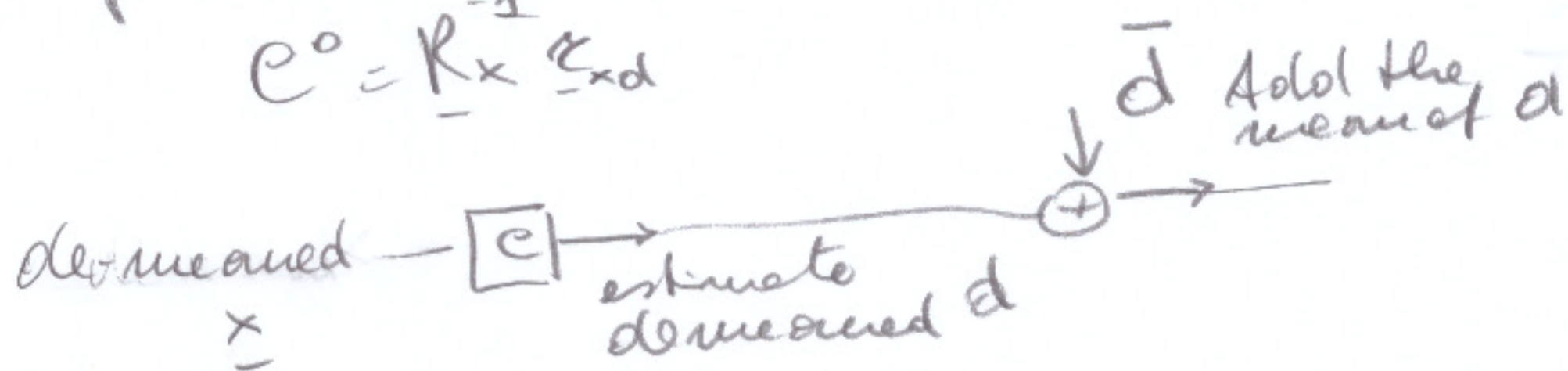
Also in considering dependence among
variables we should use covariances
 $E[(x - \mu_x)(y - \mu_y)]$ and not correlations
 $E[xy]$. In practice this issue is
avoided removing the means before
solving the problem

$$\underline{d} \leftarrow \underline{d} - E[\underline{d}]$$

$$\underline{x} \leftarrow \underline{x} - E[\underline{x}]$$

The mean of d can be added at the end after
solving the problem using correlations,
i.e. if x and d are "de-meanned"

$$\underline{c}^0 = \underline{R}_x^{-1} \underline{r}_{xd}$$



To avoid the issue of the means it is very common to generalize some linear systems into offline systems with the addition of a bias

$$y[n] = \underline{c}^T \underline{x}[n] + b = [c_1 \ c_2 \ \dots \ c_N] \begin{bmatrix} x_1[n] \\ \vdots \\ x_N[n] \end{bmatrix} + b$$

Note that an offline system is not linear because it does not satisfy the superposition principle

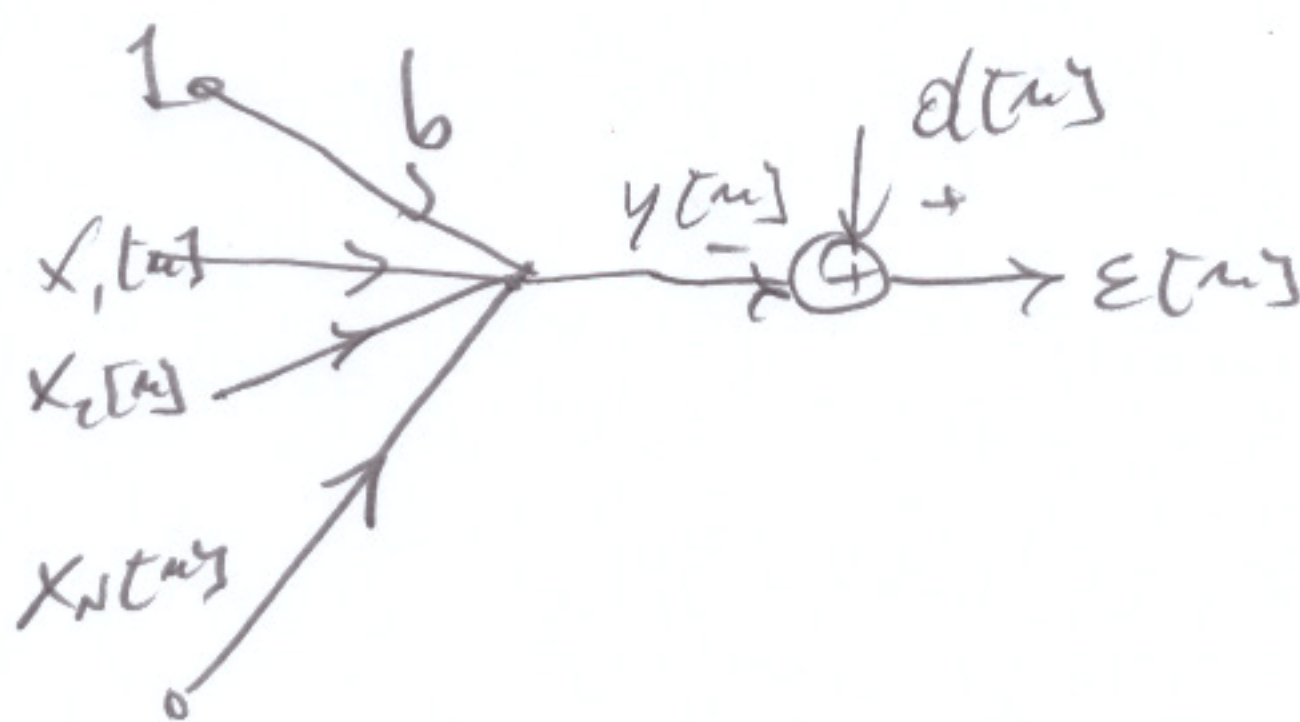
$$y_1 = c_1 x_1 + b \quad y_2 = c_2 x_2 + b$$

$$x_1 + x_2 \rightarrow c(x_1 + x_2) + b \neq y_1 + y_2$$

The constant b can be easily included in the linear system by augmenting the vector $\underline{x}[n]$ with a constant ($=1$).

$$y[n] = [c_1 \ \dots \ c_N] \begin{bmatrix} x_1[n] \\ \vdots \\ x_N[n] \end{bmatrix} + b = [b \ c_1 \ \dots \ c_N] \underbrace{\begin{bmatrix} 1 \\ x_1[n] \\ \vdots \\ x_N[n] \end{bmatrix}}_{\underline{x}_a[n]}$$

$\underline{x}_a[n]$ is an augmented vector and the problem becomes linear



$$\begin{cases} \underline{e}^o = \underline{R}_{x_a}^{-1} \underline{z}_{dx_a} \\ \underline{\Sigma}(\underline{e}^o) = \underline{\Sigma}_d - \underline{z}_{dx_a}^T \underline{R}_{x_a}^{-1} \underline{z}_{dx_a} \end{cases}$$

$$\underline{R}_{x_a} = E[\underline{x}_a[n] \underline{x}_a^T[n]] = E \left[\begin{bmatrix} 1 \\ x_1[n] \\ \vdots \\ x_N[n] \end{bmatrix} \begin{bmatrix} 1 & x_1[n] & \dots & x_N[n] \end{bmatrix} \right]$$

$$= \begin{bmatrix} 1 & E[x_1[n]] & \dots & E[x_N[n]] \\ E[x_1[n]] & & & \\ \vdots & & E[x[n]x^T[n]] & \\ E[x_N[n]] & & & \end{bmatrix}$$

↑ Note that the means are automatically included

$$\underline{z}_{dx_a} = E[d[n] \underline{x}_a[n]] = E \left[d[n] \begin{bmatrix} 1 \\ x_1[n] \\ \vdots \\ x_N[n] \end{bmatrix} \right] = \begin{bmatrix} E[d[n]] \\ E[d[n]x_1[n]] \\ \vdots \\ E[d[n]x_N[n]] \end{bmatrix} = \begin{bmatrix} E[d[n]] \\ \underline{z}_{dx} \end{bmatrix}$$

Therefore the augmented version takes into account the means implicitly and can be treated without the effect of "choosing the mean" explained above, using correlations.

Come verifica ulteriore dopo aver introdotto il bias, ovvero

$$y[n] = \underline{c}^T \underline{x}[n] + b$$

l'errore quadratico diventa

$$\mathcal{E}(\underline{c}, b) = E[(d[n] - (\underline{c}^T \underline{x}[n] + b))^2]$$

Derivando rispetto a b , abbiamo

$$\frac{\partial \mathcal{E}(\underline{c}, b)}{\partial b} = E[2(d[n] - \underline{c}^T \underline{x}[n] - b)] = 0$$

Il valore ottimo di b è

$$b^0 = E[d[n]] - \underline{c}^T E[\underline{x}[n]]$$

che sostituito nell'espressione dell'errore medio, fornisce.

$$\mathcal{E}(\underline{c}, b^0) = E[(d[n] - \underline{c}^T \underline{x}[n] - E[d[n]] + \underline{c}^T E[\underline{x}[n]])^2]$$

$$\mathcal{E}(\underline{c}, b^0) = E\left[\underbrace{(d[n] - E[d[n]])}_{\text{deviazione } d} - \underline{c}^T \underbrace{(\underline{x}[n] - E[\underline{x}[n]])}_{\text{deviazione } \underline{x}[n]}\right]^2$$

CASO SINGOLARE

MSEB

Quando la matrice $\underline{R}_x$ è singolare, è evidente che la soluzione $\underline{c}^0 = \underline{R}_x^{-1} \underline{c}_d$ non può essere calcolata immediatamente.

Ricordiamo che $\underline{R}_x = E[\underline{x}(n) \underline{x}^T(n)]$ è definita (semidefinita) positiva

Prova Una matrice $\underline{A}$ è def. positiva se $\underline{\xi}^T \underline{A} \underline{\xi} \geq 0 \quad \forall \underline{\xi}$
Nel nostro caso
$$\underline{\xi}^T E[\underline{x}(n) \underline{x}^T(n)] \underline{\xi}$$
$$= E[\underline{\xi}^T \underline{x}(n) \underline{x}^T(n) \underline{\xi}] = E[(\underline{\xi}^T \underline{x}(n))^2] \geq 0$$
$$\underline{R}_x \text{ è anche simmetrica, } \underline{R}_x^T = E[\underline{x}(n) \underline{x}^T(n)] = \underline{R}_x$$

Quindi ammette la seguente decomposizione.

$$\underline{R}_x = \underline{W} \underline{\Lambda}_x \underline{W}^T$$

autovettori autovettori autovettori
autovettori autovettori autovettori

$\underline{\Lambda}_x = \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_N)$ Gli autovalori sono tutti non negativi. Supponiamo che essi siano ordinati in ordine decrescente

$$\lambda_1 \geq \lambda_2 \geq \lambda_3 \geq \dots \geq \lambda_N$$

Gli autovettori sono "ortonormali", pertanto

$$\underline{W}^T \underline{W} = \underline{I}_N$$

$$\begin{bmatrix} \text{---} \\ \text{---} \\ \text{---} \end{bmatrix} \begin{bmatrix} \text{---} \\ \text{---} \\ \text{---} \end{bmatrix} = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}$$

($\underline{W}$ è una matrice unitaria.)

Poiché $\underline{W} \underline{W}^T \underline{W} = \underline{W}$

anche $\underline{W} \underline{W}^T = \underline{I}_N$

Nel caso singolare, alcuni autovalori sono nulli, pertanto
 la decomposizione può essere visualizzata come segue
 omettendo i zeri

$$R_x = \begin{bmatrix} \underline{W}_z & \underline{W}_{N-z} \end{bmatrix} \begin{bmatrix} \underline{\Lambda}_z & \underline{0} \\ \underline{0} & \underline{0}_{N-z} \end{bmatrix} \begin{bmatrix} \underline{W}_z^T \\ \underline{W}_{N-z}^T \end{bmatrix}$$

$\begin{bmatrix} | & \dots & | & \dots & | \end{bmatrix}$ $\begin{bmatrix} \diagup & & \\ & \square & \\ & & \diagdown \end{bmatrix}$ $\begin{bmatrix} \hline \vdots \\ \hline \vdots \\ \hline \vdots \end{bmatrix}$

autovettori dello spazio non nullo autovettori dello spazio nullo autovettori dello spazio non nullo
 autovettori dello spazio nullo

Pertanto $\underline{x}[n]$ può essere proiettato nello spazio non nullo senza perdita di informazioni



$$\underline{x}_z = \underline{W}_z^T \underline{x}$$

e lo stimatore lineare può essere applicato a $\underline{x}_z$. Ci servono $R_{x_z} = E[\underline{x}_z[n] \underline{x}_z^T[n]]$ e $\sigma_{d_{x_z}} = E[d[n] \underline{x}_z^T[n]]$
 Entrambe

$$R_{x_z} = E[\underline{x}_z[n] \underline{x}_z^T[n]] = E[\underline{W}_z^T \underline{x}[n] \underline{x}^T[n] \underline{W}_z] = \underline{W}_z^T E[\underline{x}[n] \underline{x}^T[n]] \underline{W}_z = \underline{W}_z^T R_x \underline{W}_z$$

$$= \underline{W}_z^T \underline{W}_z \underline{\Lambda}_x \underline{W}_z^T \underline{W}_z =$$

$$\underline{W}_z^T \begin{bmatrix} \underline{W}_z & \underline{W}_{N-z} \end{bmatrix} \begin{bmatrix} \underline{\Lambda}_z & \underline{0} \\ \underline{0} & \underline{0}_{N-z} \end{bmatrix} \begin{bmatrix} \underline{W}_z^T \\ \underline{W}_{N-z}^T \end{bmatrix} \underline{W}_z = \begin{bmatrix} \underline{W}_z^T \underline{W}_z & \underline{W}_z^T \underline{W}_{N-z} \\ \underline{I}_z & \underline{0}_{N-z} \end{bmatrix} \begin{bmatrix} \underline{\Lambda}_z \\ \underline{0} \end{bmatrix} \begin{bmatrix} \underline{W}_z^T \underline{W}_z \\ \underline{W}_{N-z}^T \underline{W}_z \end{bmatrix}$$

$\underline{I}_z$
 $\underline{0}_{N-z}$

$$= \underline{\Lambda}_z$$

Diagonale
 → componenti di $\underline{x}_z$ correlate

Il vettore di cross-correlazione diventa

HSE. 8

$$\begin{aligned}\underline{\zeta}_{dx_2} &= E[d[n] \underline{x}_2[n]] = E[d[n] \underline{W}_2^T \underline{x}[n]] \\ &= E[\underline{W}_2^T \underline{x}[n] d[n]] = \underline{W}_2^T E[d[n] \underline{x}[n]] \\ &= \underline{W}_2^T \underline{\zeta}_{dx}\end{aligned}$$

La soluzione per i coefficienti diventa

$$\underline{c}_2^0 = \underline{R}_{x_2}^{-1} \underline{\zeta}_{dx_2} = \underline{\Lambda}_2^{-1} \underline{W}_2^T \underline{\zeta}_{dx}$$

L'errore in condizione di ottimalità è quindi

$$E(\underline{e}_2^0) = E_d - \underline{\zeta}_{dx_2}^T \underline{R}_{x_2}^{-1} \underline{\zeta}_{dx_2} = E_d - \underline{\zeta}_{dx}^T \underline{W}_2 \underline{\Lambda}_2^{-1} \underline{W}_2^T \underline{\zeta}_{dx}$$

— 0 —

THE ORTHOGONALITY PRINCIPLE

We have derived the solution to the minimum mean squared error using the gradient. There is another interesting way of obtaining the solution using orthogonality.

Recall that the equation we solve for optimum linear estimator is

$$\underline{R}_x \underline{c}^o = \underline{r}_{dx}$$

$$E[\underline{x} \underline{x}^T] \underline{c}^o = E[\underline{D} \underline{x}]$$

that can be written as

$$E[(\underline{D} - \underline{c}^{oT} \underline{x}) \underline{x}] = \underline{0}$$

In this form (NORMAL EQUATION) we expressed the orthogonality principle.

Optimal linear mean square error estimates correspond to orthogonality between data and error.

$$E[\underline{\varepsilon}^o \underline{x}] = \underline{0}$$

We have also that the correlation between the error and the desired in optimal conditions is the minimum error

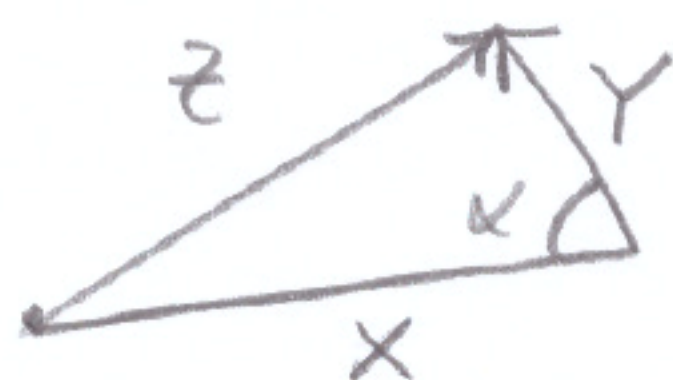
$$E[\underline{D}(\underline{D} - \underline{c}^{oT} \underline{x})] = E[\underline{D}^2] - \underline{c}^{oT} E[\underline{D} \underline{x}] = \underline{\varepsilon}^o$$

VECTOR-SPACE INTERPRETATION

USE 10

Random variables can be associated to vectors in an abstract vector space. The inner product is the covariance $E[XY]$. Orthogonality between X and Y is expressed as $E[XY] = 0$. To see what this means, consider two vectors (random variables) X and Z . The vector difference

$Y = Z - X$ is depicted as

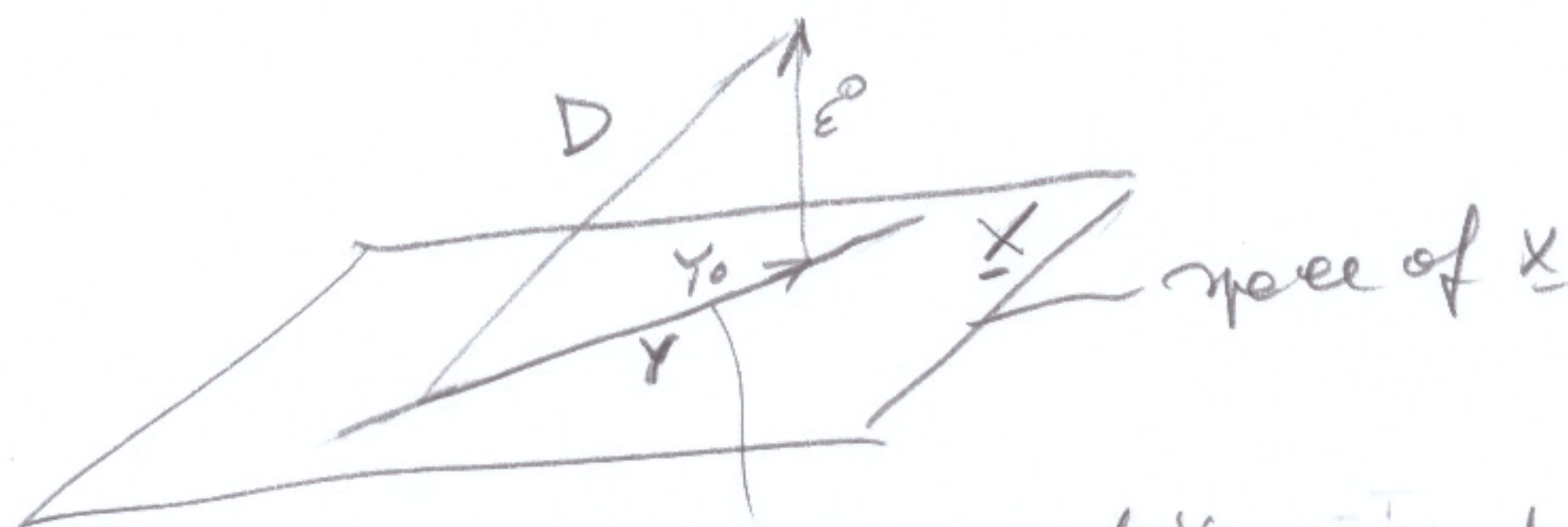
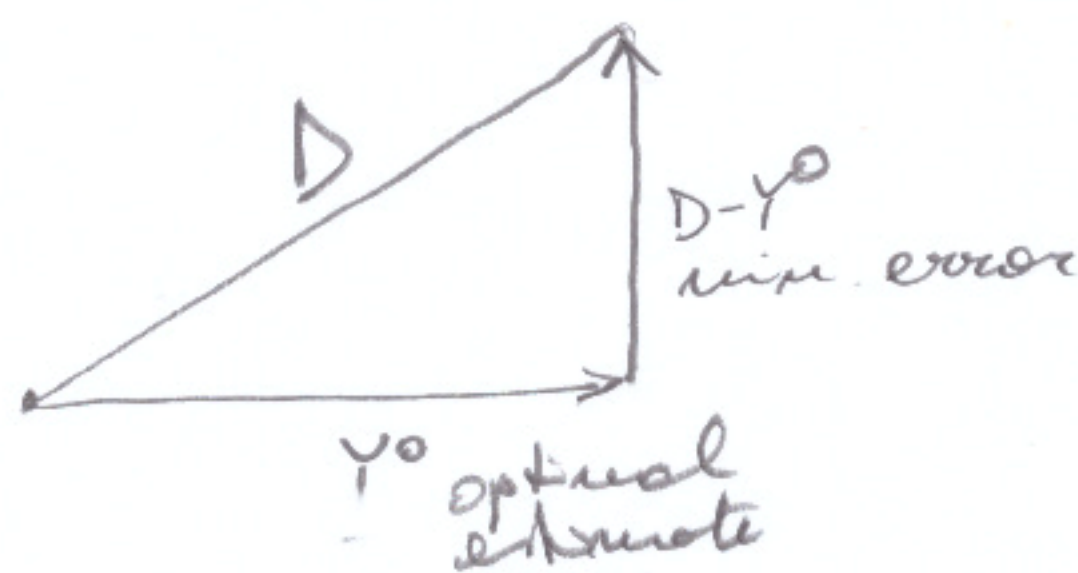
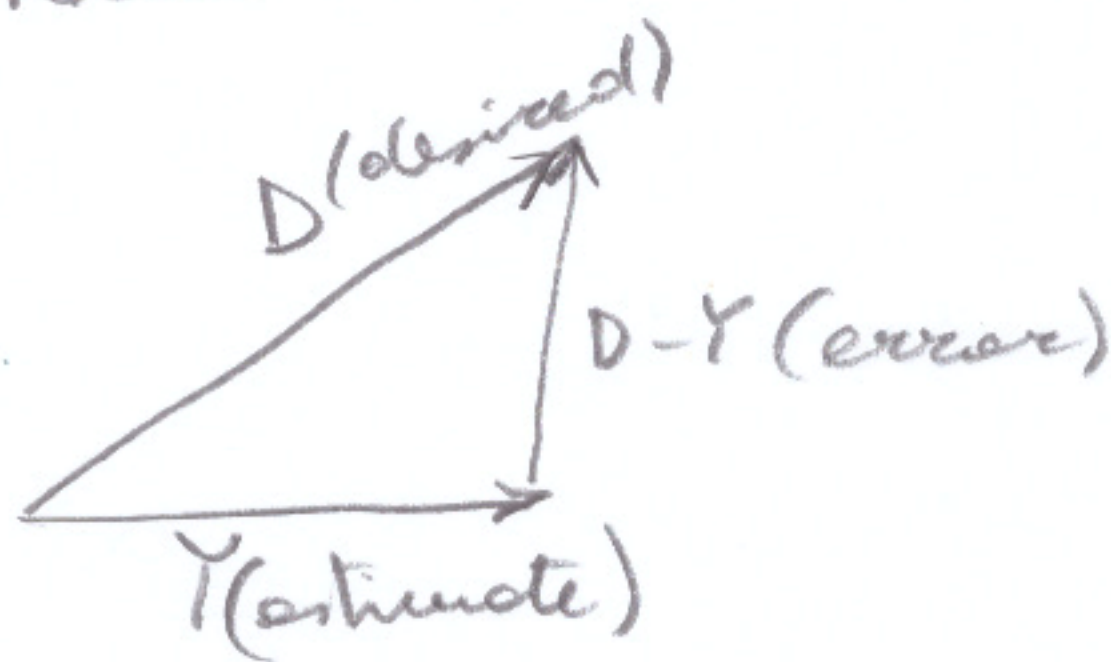


Now the Pythagorean theorem applies if $\alpha = 90^\circ$, i.e. X and Y are orthogonal.

In fact, $E[Z^2] = E[(X+Y)^2] = E[X^2] + E[Y^2] + 2E[XY]$

only if $E[XY] = 0$. The smallest difference corresponds to the vector orthogonal to X .

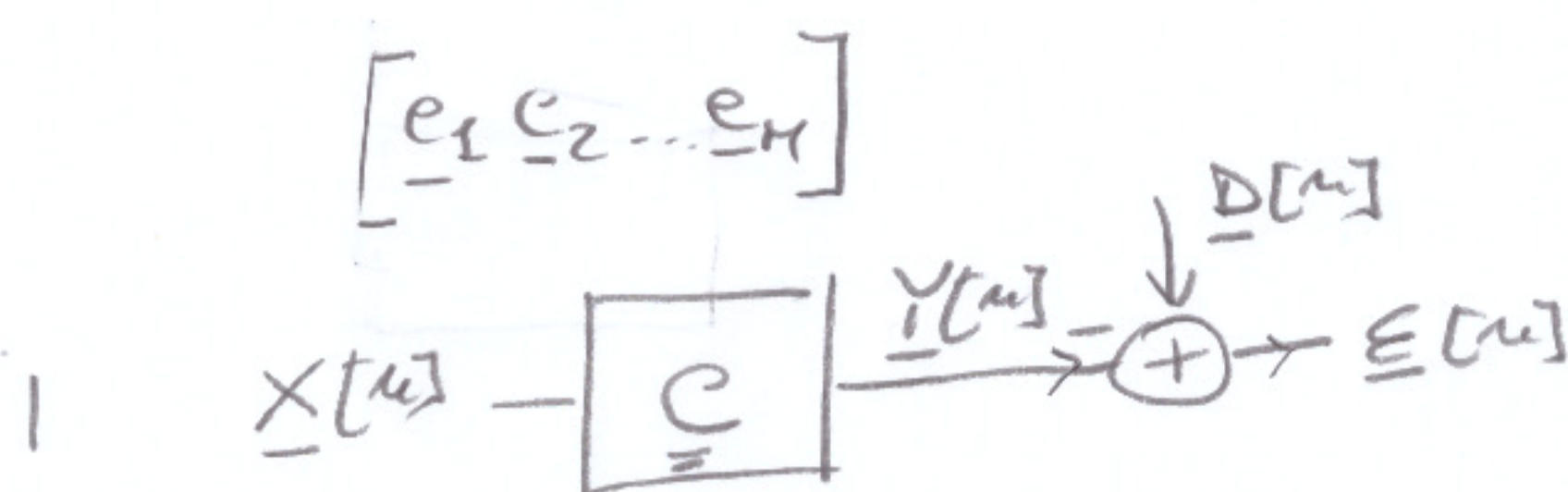
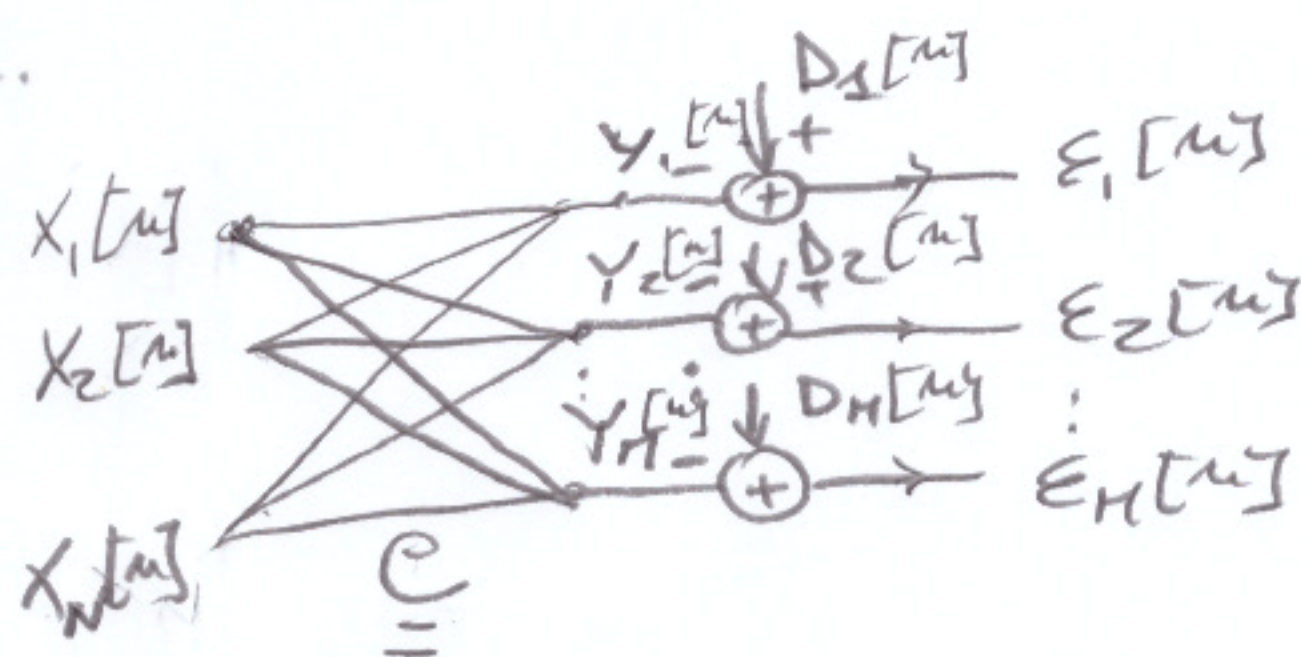
In the estimation case



linear space of Y spanned by linear functions $Y = \underline{C}^T \underline{X} \quad \forall \underline{C}$

THE MATRIX CASE (PARALLEL FILTERS)

Using the same framework presented before, we could consider a bank of M filters (linear combiners) to be optimized in parallel.



The cost function to be minimized can be taken to be the sum of all mean squared values

$$E(\underline{C}) = E[(\underline{D} - \underline{Y})^T (\underline{D} - \underline{Y})] = \sum_{i=1}^M E[(D_i - Y_i)^2] \quad (1)$$

Therefore we can optimize separately each filter to satisfy the normal equations

$$\underline{R}_x \underline{C}_i = \underline{r}_{dix} \quad i = 1, \dots, M$$

(1) Sometimes when more emphasis is given to one output with respect to the others, a more general cost could be used $E[(\underline{D} - \underline{Y})^T \underline{Q} (\underline{D} - \underline{Y})]$ where $\underline{Q}$ is symmetric and typically diagonal. We disregard this generality for now and take it to be $\underline{Q} = \underline{I}_M$.

The mean squared error for each filter ^{MSE.12} can be written in one of the equivalent forms

$$\begin{aligned}\varepsilon_i^0 &= \varepsilon_{di} - \underline{z}_{dix}^T \underline{R}_x^{-1} \underline{z}_{dix} \\ &= \varepsilon_{di} - \underline{e}_i^{0T} \underline{z}_{dix} \\ &= \varepsilon_{di} - \underline{z}_{di} y_0\end{aligned}$$

Reorganizing into matrices

$$\begin{aligned}\underline{R}_x [\underline{e}_1^0 \dots \underline{e}_M^0] &= [\underline{z}_{d1x} \underline{z}_{d2x} \dots \underline{z}_{dMx}] \\ E[\underline{x} \underline{x}^T \underline{C}^0] &= E[\underline{D}_1 \underline{x} \underline{D}_2 \underline{x} \dots \underline{D}_M \underline{x}] \\ E[\underline{x} \underline{x}^T \underline{C}^0] &= E[\underline{x} \underline{D}^T] \quad \underline{R}_{xD} \triangleq E[\underline{x} \underline{D}^T]\end{aligned}$$

$$\boxed{\underline{R}_x \underline{C}^0 = \underline{R}_{xD}} \Rightarrow \underline{C}^0 = \underline{R}_x^{-1} \underline{R}_{xD}$$

The total squared error in optimal condition is

$$\varepsilon^0 = \sum_{i=1}^M \varepsilon_i^0 = \sum_{i=1}^M \varepsilon_{di} - \sum_{i=1}^M \underline{z}_{dix}^T \underline{R}_x^{-1} \underline{z}_{dix}$$

Note that, even if the optimization for each filter is computed separately, the errors may be dependent also because $\underline{D}_1, \underline{D}_2, \dots, \underline{D}_M$ may be dependent and have a correlation $\underline{R}_D = E[\underline{D} \underline{D}^T]$

Therefore it may be useful to consider the error correlation matrix

$$\underline{R}_\varepsilon = E[\underline{\varepsilon} \underline{\varepsilon}^T] = E[(\underline{D} - \underline{C}^T \underline{x})(\underline{D} - \underline{C}^T \underline{x})^T]$$

that in optimal conditions is

MSE.13

$$\underline{R}_{\underline{\varepsilon}^0} = E[\underline{\varepsilon}^0 \underline{\varepsilon}^{0T}] = E[(\underline{D} - \underline{C}^{0T} \underline{x})(\underline{D} - \underline{C}^{0T} \underline{x})^T]$$

$$= E[\underline{D} \underline{D}^T] + \underline{C}^{0T} E[\underline{x} \underline{x}^T] \underline{C}^0 - \underline{C}^{0T} E[\underline{x} \underline{D}^T] - E[\underline{D} \underline{x}^T] \underline{C}^0$$

$$= \underline{R}_D + \underline{R}_{XD} \underline{R}_X^{-1} \underline{R}_X \underline{R}_{XD} - \underline{R}_{XD} \underline{R}_X^{-1} \underline{R}_{XD} - \underline{R}_{XD} \underline{R}_X^{-1} \underline{R}_{XD}$$

$$= \underline{R}_D - \underline{R}_{XD} \underline{R}_X^{-1} \underline{R}_{XD}$$

$$= \underline{R}_D - \underline{C}^{0T} \underline{R}_{XD}$$

It is easy to see that the total squared error can be written as

$$\begin{aligned} \varepsilon^0 &= \text{tr}(\underline{R}_{\underline{\varepsilon}^0}) = \text{tr}(\underline{R}_D - \underline{R}_{XD} \underline{R}_X^{-1} \underline{R}_{XD}) \\ &= \text{tr}(\underline{R}_D - \underline{C}^{0T} \underline{R}_{XD}) \end{aligned}$$

ADDING LINEAR CONSTRAINTS

If we have constraints on the filter that can be expressed linearly in the coefficients, it is possible to find closed form solution.

ONE LINEAR CONSTRAINT

$$\begin{cases} \min_{\underline{c}} E(\underline{c}) \\ \underline{s}^T \underline{c} = \mu \end{cases}$$

$$\left(\begin{array}{l} \text{a linear constraint} \\ \text{may be } \sum_{i=1}^N c_i = 1? \\ \underline{c}^T \underline{e} = 1 \end{array} \right)$$

The Lagrange function is then

$$L(\underline{c}, \lambda) = E[(d - \underline{c}^T \underline{x})^2] + \lambda (\underline{c}^T \underline{s} - \mu)$$

$$= E[d^2] - 2 \underline{c}^T \underline{c}_{dx} + \underline{c}^T \underline{R}_x \underline{c} + \lambda (\underline{c}^T \underline{s} - \mu)$$

The gradient equation is

$$\nabla_{\underline{c}} L(\underline{c}, \lambda) = -2 \underline{c}_{dx} + 2 \underline{R}_x \underline{c} + \lambda \underline{s} = \underline{0} \quad (a)$$

Multiply by $\underline{R}_x^{-1}$ (on the left)

$$-2 \underline{R}_x^{-1} \underline{c}_{dx} + 2 \underline{c} + \lambda \underline{R}_x^{-1} \underline{s} = \underline{0}$$

Multiply by $\underline{s}^T$ both sides.

$$-2 \underline{s}^T \underline{R}_x^{-1} \underline{c}_{dx} - 2 \underbrace{\underline{s}^T \underline{c}}_{\mu} + \lambda \underline{s}^T \underline{R}_x^{-1} \underline{s} = 0$$

$$\lambda = 2 \frac{\underline{s}^T \underline{R}_x^{-1} \underline{c} - \underline{s}^T \underline{c}}{\underline{s}^T \underline{R}_x^{-1} \underline{s}}$$

that substituted back in (a) gives:

$$-\cancel{2} \underline{c}_{dx} + \cancel{2} \underline{R}_x \underline{c} + \cancel{2} \frac{\underline{s}^T \underline{R}_x^{-1} \underline{c} - \underline{s}^T \underline{c}}{\underline{s}^T \underline{R}_x^{-1} \underline{s}} \underline{s} = 0$$

$$\boxed{\underline{c}^0 = \underline{R}_x^{-1} \underline{c}_{dx} - \frac{\underline{s}^T \underline{R}_x^{-1} \underline{s} - \mu}{\underline{s}^T \underline{R}_x^{-1} \underline{s}} \underline{R}_x^{-1} \underline{s}}$$

M LINEAR CONSTRAINTS

MSE.15

$$\left\{ \begin{array}{l} \min_{\underline{c}} \mathcal{E}(\underline{c}) \\ \underline{S}^T \underline{c} = \underline{\mu} \\ \begin{array}{c} N \\ H \end{array} \quad \begin{array}{c} N \\ H \end{array} \end{array} \right.$$

Linear constraints could be introduced, for example, to impose response to a set of given frequencies.

$$L(\underline{c}, \underline{\lambda}) = \underline{c}^T \underline{R}_x \underline{c} - 2 \underline{\tau}_{dx}^T \underline{c} + \mathcal{E}_d + (\underline{c}^T \underline{S} - \underline{\mu}^T) \underline{\lambda}$$

↑
M Lagrange multipliers

$$\nabla_{\underline{c}} L(\underline{c}, \underline{\lambda}) = 2 \underline{R}_x \underline{c} - 2 \underline{\tau}_{dx} + \underline{S} \underline{\lambda} = \underline{0} \quad (*)$$

Multiply by $\underline{R}_x^{-1}$ on the left.

$$2 \underline{R}_x^{-1} \underline{R}_x \underline{c} - 2 \underline{R}_x^{-1} \underline{\tau}_{dx} + \underline{R}_x^{-1} \underline{S} \underline{\lambda} = \underline{0}$$

Multiply both sides by $\underline{S}^T$ (on the left).

$$2 \underline{S}^T \underline{c} - 2 \underline{S}^T \underline{R}_x^{-1} \underline{\tau}_{dx} + \underline{S}^T \underline{R}_x^{-1} \underline{S} \underline{\lambda} = \underline{0}$$

$$\underline{\lambda} = 2 (\underline{S}^T \underline{R}_x^{-1} \underline{S})^{-1} (\underline{S}^T \underline{R}_x^{-1} \underline{\tau}_{dx} - \underline{S}^T \underline{c})$$

$\underline{\mu}$

Substituting into (*)

$$2 \underline{R}_x \underline{c} - 2 \underline{\tau}_{dx} + \underline{S} \left((\underline{S}^T \underline{R}_x^{-1} \underline{S})^{-1} (\underline{S}^T \underline{R}_x^{-1} \underline{\tau}_{dx} - \underline{\mu}) \right) = \underline{0}$$

$$\underline{c} = \underline{R}_x^{-1} \underline{\tau}_{dx} - \underline{R}_x^{-1} \underline{S} (\underline{S}^T \underline{R}_x^{-1} \underline{S})^{-1} (\underline{S}^T \underline{R}_x^{-1} \underline{\tau}_{dx} - \underline{\mu})$$

This expression clearly requires the invertibility of $\underline{S}^T \underline{R}_x \underline{S}$ and must be checked case-by-case.

THE CANONICAL FORM

It is often useful to write the error in the so-called canonical form

$$\boxed{\mathcal{E}(\underline{c}) = \mathcal{E}(\underline{c}^0) + (\underline{c} - \underline{c}^0)^T \underline{R}_x (\underline{c} - \underline{c}^0)}$$

Where $\underline{c}^0 = \underline{R}_x^{-1} \underline{z}_{dx}$ and $\mathcal{E}(\underline{c}^0) = \mathcal{E}_d - \underline{z}_{dx}^T \underline{R}_x^{-1} \underline{z}_{dx}$

The canonical form can be verified by direct substitution:

$$\begin{aligned} \mathcal{E}(\underline{c}) &= \mathcal{E}_d - \underline{z}_{dx}^T \underline{R}_x^{-1} \underline{z}_{dx} + \underline{c}^T \underline{R}_x \underline{c} + \underline{c}^{0T} \underline{R}_x \underline{c}^0 - \underline{c}^T \underline{R}_x \underline{c}^0 - \underline{c}^{0T} \underline{R}_x \underline{c} \\ &= \mathcal{E}_d - \cancel{\underline{z}_{dx}^T \underline{R}_x^{-1} \underline{z}_{dx}} + \underline{c}^T \underline{R}_x \underline{c} + \cancel{\underline{z}_{dx}^T \underline{R}_x^{-1} \underline{R}_x \underline{R}_x^{-1} \underline{z}_{dx}} - \cancel{\underline{c}^T \underline{R}_x \underline{R}_x^{-1} \underline{z}_{dx}} \\ &\quad - \cancel{\underline{z}_{dx}^T \underline{R}_x^{-1} \underline{R}_x \underline{R}_x^{-1} \underline{c}} \\ &= \mathcal{E}_d + \underline{c}^T \underline{R}_x \underline{c} - 2 \underline{z}_{dx}^T \underline{c} \quad \square \end{aligned}$$