

Modern Digital Communications: A Hands-On Approach

MMSE Estimation and LS Approximation

Dr. Nicolae Chiurtu

- Course material of Prof. Bixio Rimoldi -

Last revision: Nov. 27, 2023

Minimum Mean Squared Error Estimation

In this note, we study the concept of Minimum Mean Squared Error (MMSE) Estimation. We start with a general description of the problem, then focus on the special case of MMSE estimation for jointly Gaussian random vectors — which is what we need in order to estimate the OFDM channel.

Subsequently we study Least Squares (LS) approximations, which will also be used in the assignments.

Estimation

We are interested in estimating the realization x of the random variable $\mathbf{X} \in \mathbb{C}^m$.

For this, we have access to the realization $\mathbf{y}$ of $\mathbf{Y} \in \mathbb{C}^n$. (Here we use small letters to denote the realization of random variables). Unless $\mathbf{X}$ and $\mathbf{Y}$ are independent, knowing $\mathbf{y}$ should help estimating x .

An *estimator* of $\mathbf{X} \in \mathbb{C}^m$ based on $\mathbf{Y} \in \mathbb{C}^n$ can be an arbitrary function $\hat{\mathbf{X}}: \mathbb{C}^n \rightarrow \mathbb{C}^m$. We then let $\hat{x} = \hat{\mathbf{X}}(\mathbf{y})$ be the estimate of x based on the observed $\mathbf{y}$. Clearly this estimate may or may not be a good one.

Thus we need a way to measure the performance of an estimator, and hopefully we can find an estimator which is optimal under that performance criterion.

MMSE Estimator

A popular choice for the fidelity criterion is the *mean squared error*

$$\text{MSE}(\hat{\mathbf{X}}) := E \left[\|\mathbf{X} - \hat{\mathbf{X}}(\mathbf{Y})\|^2 \right]$$

where the expectation is over the joint distribution of $\mathbf{X}$ and $\mathbf{Y}$.

An estimator that minimizes the mean squared error is called a *minimum mean squared error (MMSE) estimator* of $\mathbf{X}$ given $\mathbf{Y}$.

We denote by $\hat{\mathbf{X}}_{\text{MMSE}}(\mathbf{y})$ such an estimator. In other words, for any estimator $\hat{\mathbf{X}}(\mathbf{y})$,

$$E \left[\|\mathbf{X} - \hat{\mathbf{X}}(\mathbf{Y})\|^2 \right] \geq E \left[\|\mathbf{X} - \hat{\mathbf{X}}_{\text{MMSE}}(\mathbf{Y})\|^2 \right].$$

The MMSE estimator of $\mathbf{X}$ given $\mathbf{Y}$ turns out to be *unique* and equal to the *conditional expectation* of $\mathbf{X}$ given $\mathbf{Y}$, i.e.,

$$\hat{\mathbf{X}}_{\text{MMSE}}(\mathbf{y}) = E[\mathbf{X} | \mathbf{Y} = \mathbf{y}].$$

Derivation of the MMSE Estimator

We are looking for an estimator $\hat{\mathbf{X}}: \mathbb{C}^n \rightarrow \mathbb{C}^m$ such that

$$\text{MSE}(\hat{\mathbf{X}}) = E \left[\|\mathbf{X} - \hat{\mathbf{X}}(\mathbf{Y})\|^2 \right],$$

is minimized.

We are done if we prove that $\hat{\mathbf{X}}_{\text{MMSE}}(\mathbf{y})$ minimizes

$$E_y \left[\|\mathbf{X} - \hat{\mathbf{X}}(\mathbf{Y})\|^2 \right] \triangleq E \left[\|\mathbf{X} - \hat{\mathbf{X}}(\mathbf{Y})\|^2 \middle| \mathbf{Y} = \mathbf{y} \right],$$

where E_y means that we are taking the expectation conditioning on $\mathbf{Y} = \mathbf{y}$.

We have

$$E_y \left[\|\mathbf{X} - \hat{\mathbf{X}}(\mathbf{Y})\|^2 \right] = E_y \left[\|\mathbf{X}\|^2 \right] + \|\hat{\mathbf{X}}(\mathbf{y})\|^2 - 2\Re \left\{ E_y \left[\mathbf{X}^\dagger \right] \hat{\mathbf{X}}(\mathbf{y}) \right\}.$$

The first term on the RHS of the equality does not depend on $\hat{\mathbf{X}}$. If we substitute it for another term that does not depend on $\hat{\mathbf{X}}$, then the minimizing $\hat{\mathbf{X}}$ remains the same.

We choose to minimize

$$\|E_y[\mathbf{X}]\|^2 + \|\hat{\mathbf{X}}(\mathbf{y})\|^2 - 2\Re \left\{ E_y[\mathbf{X}^\dagger] \hat{\mathbf{X}}(\mathbf{y}) \right\},$$

which is the quadratic form

$$\|E_y[\mathbf{X}] - \hat{\mathbf{X}}(\mathbf{y})\|^2.$$

The above is non-negative, and it achieves its minimum iff $\hat{\mathbf{X}}(\mathbf{y}) = E_y[\mathbf{X}]$.

□

MMSE Estimator for Jointly Gaussian Random Vectors

In general, finding an expression for $E[\mathbf{X}|\mathbf{Y} = \mathbf{y}]$ is not a trivial task.

However, when $\mathbf{X}$ and $\mathbf{Y}$ are jointly Gaussian random vectors, the MMSE estimator has a simple linear form.

In the rest of this lecture we restrict our discussion to jointly Gaussian random vectors $\mathbf{X}$ and $\mathbf{Y}$.

First we assume that $\mathbf{X}$ and $\mathbf{Y}$ are zero-mean vectors. Later on, we will take non-zero means into account.

A Fact About Jointly Gaussian Random Vectors

If $\mathbf{X} \in \mathbb{C}^m$ and $\mathbf{Y} \in \mathbb{C}^n$ are zero-mean jointly Gaussian random vectors, we can always write

$$\mathbf{X} = A\mathbf{Y} + \mathbf{Z}$$

for some matrix $A \in \mathbb{C}^{m \times n}$ (that we will determine shortly) and a zero-mean Gaussian vector $\mathbf{Z} \sim \mathcal{N}_{\mathcal{C}}(\mathbf{0}, K_{\mathbf{Z}})$ that is independent of $\mathbf{Y}$.

Proof. For any matrix $A \in \mathbb{C}^{m \times n}$,

$$\mathbf{Z} := \mathbf{X} - A\mathbf{Y}$$

is a (zero-mean) complex Gaussian vector. (By definition, a linear combination of jointly Gaussian vectors leads to a Gaussian vector.)

Moreover,

$$E[\mathbf{Z}\mathbf{Y}^\dagger] = E[\mathbf{X}\mathbf{Y}^\dagger] - A E[\mathbf{Y}\mathbf{Y}^\dagger] = K_{\mathbf{X}\mathbf{Y}} - A K_{\mathbf{Y}}$$

Hence, if we take

$$A = K_{\mathbf{X}\mathbf{Y}}K_{\mathbf{Y}}^{-1},$$

we will have

$$E [\mathbf{Z}\mathbf{Y}^\dagger] = 0,$$

i.e., with this specific choice of A , $\mathbf{Z}$ is independent of $\mathbf{Y}$. □

Note that with this choice of A ,

$$\begin{aligned} K_{\mathbf{Z}} &= E [(\mathbf{X} - A\mathbf{Y})(\mathbf{X} - A\mathbf{Y})^\dagger] \\ &= K_{\mathbf{X}} - AK_{\mathbf{Y}\mathbf{X}} - (K_{\mathbf{X}\mathbf{Y}} - AK_{\mathbf{Y}})A^\dagger \\ &= K_{\mathbf{X}} - AK_{\mathbf{Y}\mathbf{X}} \\ &= K_{\mathbf{X}} - K_{\mathbf{X}\mathbf{Y}}K_{\mathbf{Y}}^{-1}K_{\mathbf{X}\mathbf{Y}}^\dagger, \end{aligned}$$

where in the third line we used the fact that $AK_{\mathbf{Y}} = K_{\mathbf{X}\mathbf{Y}}$.

MMSE Estimator for Jointly Gaussian Vectors

Let $\mathbf{X}$ and $\mathbf{Y}$ be jointly Gaussian. Then

$$\hat{\mathbf{X}}_{\text{MMSE}}(\mathbf{y}) = K_{\mathbf{XY}} K_{\mathbf{Y}}^{-1} \mathbf{y}.$$

Moreover, the minimum mean squared error equals

$$E \left[\|\mathbf{X} - \hat{\mathbf{X}}_{\text{MMSE}}(\mathbf{Y})\|^2 \right] = \text{trace}(K_{\mathbf{X}} - K_{\mathbf{XY}} K_{\mathbf{Y}}^{-1} K_{\mathbf{XY}}^{\dagger}).$$

Proof. Write

$$\mathbf{X} = A\mathbf{Y} + \mathbf{Z} \quad \text{with} \quad A = K_{\mathbf{X}\mathbf{Y}}K_{\mathbf{Y}}^{-1},$$

which makes $\mathbf{Z}$ independent of $\mathbf{Y}$.

Now

$$\begin{aligned}\hat{\mathbf{X}}_{\text{MMSE}}(\mathbf{y}) &= E[\mathbf{X}|\mathbf{Y} = \mathbf{y}] \\ &= E[A\mathbf{Y} + \mathbf{Z}|\mathbf{Y} = \mathbf{y}] \\ &= E[A\mathbf{y} + \mathbf{Z}] = A\mathbf{y}.\end{aligned}$$

The minimum mean squared error is

$$\begin{aligned}E\left[\|\mathbf{X} - \hat{\mathbf{X}}_{\text{MMSE}}(\mathbf{Y})\|^2\right] &= E\left[\|(A\mathbf{Y} + \mathbf{Z}) - A\mathbf{Y}\|^2\right] \\ &= E\left[\|\mathbf{Z}\|^2\right] = \text{trace}(K_{\mathbf{Z}}) = \text{trace}(K_{\mathbf{X}} - K_{\mathbf{X}\mathbf{Y}}K_{\mathbf{Y}}^{-1}K_{\mathbf{X}\mathbf{Y}}^{\dagger}).\end{aligned}$$

□

Channel Coefficients Estimation

We have seen that by using OFDM we ‘forge’ parallel channels described in matrix form as

$$\mathbf{Y}^{(m)} = D\mathbf{A}^{(m)} + \mathbf{Z}^{(m)}$$

where D is the diagonal matrix of channel coefficients. How to estimate the matrix D ?

For certain values of m , we substitute $\mathbf{A}^{(m)}$ with an N -tuple $\mathbf{S}$ known to the receiver. Then, dropping the superscript (m) for notational convenience,

$$\mathbf{Y} = D\mathbf{S} + \mathbf{Z}$$

where D is diagonal. The same result is obtained by

$$\mathbf{Y} = \mathbf{S}\mathbf{D} + \mathbf{Z}$$

where $\mathbf{S}$ is the diagonal matrix that has $\mathbf{S}$ as its diagonal elements and $\mathbf{D}$ is the N -tuple consisting of the diagonal elements of D .

Assume that $\mathbf{D}$ is zero-mean, Gaussian, and independent of $\mathbf{Z}$.

Then $\mathbf{Y}$ and $\mathbf{D}$ are jointly Gaussian. To see this, write

$$\begin{pmatrix} \mathbf{Y} \\ \mathbf{D} \end{pmatrix} = \begin{pmatrix} S & I \\ I & 0 \end{pmatrix} \begin{pmatrix} \mathbf{D} \\ \mathbf{Z} \end{pmatrix}.$$

This shows that the vector on the left hand side is a linear transformation of a Gaussian random vector.

Hence, the minimum mean squared error (MMSE) estimate of $\mathbf{D}$ based on the observable $\mathbf{Y}$ is

$$\hat{\mathbf{D}} = K_{D\mathbf{Y}} K_{\mathbf{Y}}^{-1} \mathbf{Y},$$

where

$$\begin{aligned} K_{D\mathbf{Y}} &:= E[\mathbf{D}\mathbf{Y}^\dagger] = K_D S^\dagger \quad \text{and} \\ K_{\mathbf{Y}} &:= E[\mathbf{Y}\mathbf{Y}^\dagger] = S K_D S^\dagger + K_{\mathbf{Z}} \end{aligned}$$

are covariance matrices and $\mathbf{D}$ and $\mathbf{Y}$ are zero-mean.

Both $K_{\mathbf{Y}}$ and $K_{D\mathbf{Y}}$ depend on K_D . The question is how to find K_D .

K_D from the Channel Model

To determine K_D , we make a channel model.

A reasonable assumption for a wireless channel is the multipath channel, modeled as

$$h(t) = \sum_{l=0}^{M-1} \alpha_l \delta(t - \tau_l),$$

where α_l and τ_l are the l -th path strength and delay, respectively.

Since α_l and α_k , $l \neq k$, describe the reflection on different obstacles/surfaces, it is reasonable to assume that they are uncorrelated.

For this channel,

$$h_{\mathcal{F}}(f) = \int h(t) e^{-j2\pi f t} dt = \sum_l \alpha_l \int \delta(t - \tau_l) e^{-j2\pi f t} dt = \sum_l \alpha_l e^{-j2\pi f \tau_l}.$$

Assuming that the DFT-length N is even (usually a power of 2), $\mathbf{D} = (\lambda_0, \dots, \lambda_{N-1})$ with

$$\lambda_i = \begin{cases} h_{\mathcal{F}}\left(\frac{i}{NT_s}\right), & i = 0, \dots, \frac{N}{2} - 1 \\ h_{\mathcal{F}}\left(\frac{i-N}{NT_s}\right), & i = \frac{N}{2}, \dots, N - 1 \end{cases},$$

where

$$h_{\mathcal{F}}\left(\frac{i}{NT_s}\right) = \sum_l \alpha_l e^{-j2\pi \frac{i}{NT_s} \tau_l},$$

and T_s is both the symbol interval and the sampling interval.

If we define

$$[i] = \begin{cases} i, & i = 0, \dots, \frac{N}{2} - 1 \\ i - N, & i = \frac{N}{2}, \dots, N - 1 \end{cases},$$

then we can write

$$\lambda_i = h_{\mathcal{F}}\left(\frac{[i]}{NT_s}\right).$$

The (i, k) entry of the covariance matrix K_D is

$$\begin{aligned} (K_D)_{i,k} &= E [\lambda_i \lambda_k^*] = E \left[\sum_l \alpha_l e^{-j2\pi \frac{[i]}{NT_s} \tau_l} \sum_v \alpha_v^* e^{j2\pi \frac{[k]}{NT_s} \tau_v} \right] \\ &= \sum_l \sum_v e^{-j2\pi \frac{[i]}{NT_s} \tau_l} E [\alpha_l \alpha_v^*] e^{j2\pi \frac{[k]}{NT_s} \tau_v} \end{aligned} \quad (1)$$

The above expression seems complicated but it is actually not. In fact, K_D is the product of three matrices. To see this, notice that if A, B, C are matrices such that the product ABC is well defined, then

$$(ABC)_{i,k} = \sum_l \sum_v A_{i,l} B_{l,v} C_{v,k},$$

which is exactly the right-hand side of (1) with

$$A_{i,l} = e^{-j2\pi \frac{[i]}{NT_s} \tau_l}, \quad B_{l,v} = E [\alpha_l \alpha_v^*], \quad C = A^\dagger.$$

From K_D , we determine K_Y and K_{DY} as described earlier.

Least Squares (LS) Approximation

An MMSE estimator requires the knowledge of the statistics.

Suppose that all we know is that the observable $\mathbf{y} \in \mathbb{C}^n$ is obtained from $\boldsymbol{\lambda} \in \mathbb{C}^n$ according to

$$\mathbf{y} = S\boldsymbol{\lambda} + \mathbf{z},$$

where S is a diagonal matrix with non-vanishing diagonal elements, and $\mathbf{z} \in \mathbb{C}^n$ is a noise vector. None of the statistics are known.

We would like to find the $\hat{\boldsymbol{\lambda}} \in \mathbb{C}^n$ for which $\|\mathbf{y} - S\boldsymbol{\lambda}\|^2$ is minimized over all $\boldsymbol{\lambda} \in \mathbb{C}^n$.

By writing

$$\|\mathbf{y} - S\boldsymbol{\lambda}\|^2 = \sum_{i=0}^{n-1} |y_i - S_i \lambda_i|^2,$$

it is clear that $\hat{\lambda}_i$ needs to be chosen to minimize $|y_i - S_i \lambda_i|^2$.

Clearly the minimum is obtained when

$$\hat{\lambda}_i = \frac{y_i}{S_i},$$

in which case $\|\mathbf{y} - S\boldsymbol{\lambda}\|^2 = 0$.

This is a special case that we have worked out since we will need it in the assignment.

More generally, suppose that

$$\mathbf{y} = S\boldsymbol{\lambda} + \mathbf{z},$$

where $\mathbf{y}$ and $\mathbf{z}$ are in $\mathbb{C}^n$, $\boldsymbol{\lambda} \in \mathbb{C}^m$, and $S \in \mathbb{C}^{n \times m}$ is a general full-rank matrix (not necessarily diagonal).

As before, we are seeking the $\hat{\boldsymbol{\lambda}}$ that minimizes

$$\|\mathbf{y} - S\boldsymbol{\lambda}\|^2$$

over all $\boldsymbol{\lambda} \in \mathbb{C}^m$.

If $n \leq m$, the system $\mathbf{y} = S\boldsymbol{\lambda}$ can always be solved. Hence we can always find a $\boldsymbol{\lambda}$ for which $\|\mathbf{y} - S\boldsymbol{\lambda}\|^2 = 0$.

So we focus on the case where $n > m$ (the system is overdetermined).

We may reformulate the problem as follows. The observable $\mathbf{y}$ is an element of the inner-product space $\mathcal{U} = \mathbb{C}^n$ and let $\mathcal{V}$ be the subspace spanned by the columns of S . We are seeking the vector $\hat{\mathbf{y}} \in \mathcal{V}$ that minimizes

$$\|\mathbf{y} - \hat{\mathbf{y}}\|^2.$$

The projection theorem (see PDC) tells us that $\hat{\mathbf{y}}$ is the projection of $\mathbf{y}$ into $\mathcal{V}$. It has the property that the error vector $\mathbf{y} - \hat{\mathbf{y}}$ is orthogonal to every element of $\mathcal{V}$. In particular, it is orthogonal to the columns of S . Hence

$$\langle \mathbf{y} - S\hat{\boldsymbol{\lambda}}, \mathbf{S}_i \rangle = 0, \quad i = 1, \dots, m$$

where $\mathbf{S}_i$ is the i -th column of S .

Hence

$$\langle \mathbf{y}, \mathbf{S}_i \rangle = \langle S\hat{\boldsymbol{\lambda}}, \mathbf{S}_i \rangle, \quad i = 1, \dots, m. \quad (2)$$

Equivalently,

$$\mathbf{S}_i^\dagger \mathbf{y} = \mathbf{S}_i^\dagger S \hat{\boldsymbol{\lambda}} \quad i = 1, \dots, m.$$

In matrix form, we obtain the so-called *normal equations*:

$$S^\dagger \mathbf{y} = S^\dagger S \hat{\boldsymbol{\lambda}}.$$

Since S has full rank, $S^\dagger S$ is nonsingular. We prove this by arguing that if $S^\dagger S$ is singular, then S does not have full rank. Indeed, if $S^\dagger S$ is singular, there exists a nonzero vector $\mathbf{u}$ such that $S^\dagger S \mathbf{u} = 0$, which implies that $\mathbf{u}^\dagger S^\dagger S \mathbf{u} = \|S \mathbf{u}\|^2 = 0$, so that $S \mathbf{u} = 0$. This means that the columns of S are linearly dependent, i.e., S is not full rank — a contradiction.

Solving for $\hat{\boldsymbol{\lambda}}$ yields the least-squares approximation of $\boldsymbol{\lambda}$:

$$\hat{\boldsymbol{\lambda}} = (S^\dagger S)^{-1} S^\dagger \mathbf{y}.$$

To reconstruct the above formula from memory, it suffices to check the dimensions. The matrix that estimates λ from y must be of dimension $m \times n$. If we analyze $(S^\dagger S)^{-1} S^\dagger$ from right to left, we see that we don't have other choices (unless we make the expression more complicated). In particular, $S^\dagger$ has the right dimension to multiply y and notice that $(S^\dagger S)^{-1}$ is $m \times m$, hence it is a valid matrix to multiply the $m \times n$ matrix $S^\dagger$, whereas $(SS^\dagger)^{-1}$ would not fit as it is $n \times n$.

To summarize, recall the fundamental difference between the MMSE estimate and the LS approximation. On the one hand, the MMSE setup assumes two random vectors, the unobserved $\mathbf{X}$ and the observed $\mathbf{Y}$, and the objective is to estimate $\mathbf{X}$ by means of $\mathbf{Y}$, where the measure of performance is $E[\|\mathbf{X} - \hat{\mathbf{X}}(\mathbf{Y})\|^2]$.

On the other hand, the LS approximation is about adjusting the parameters of a model function to best fit the observed data. (See the example that follows.) It is also called data fitting or regression analysis. (See e.g. Least Squares in Wikipedia).

Example: Suppose that the model function is $y(x) = \alpha_1 + \alpha_2 x + \alpha_3 x^2$, with unknown parameters $\alpha_1, \alpha_2, \alpha_3$. We seek to find the parameters so that the model fits the noisy observations

$$y_i = y(x_i) + z_i, \quad i = 1, 2, \dots, n.$$

Letting $\mathbf{y} = (y_1, y_2, \dots, y_n)^T$, $\mathbf{x} = (x_1, x_2, \dots, x_n)^T$, $\boldsymbol{\alpha} = (\alpha_1, \alpha_2, \alpha_3)^T$, and $\mathbf{z} = (z_1, z_2, \dots, z_n)^T$, the observations are of the form

$$\mathbf{y} = S\boldsymbol{\alpha} + \mathbf{z},$$

where

$$S = \begin{pmatrix} 1 & x_1 & x_1^2 \\ 1 & x_2 & x_2^2 \\ \vdots & \vdots & \vdots \\ 1 & x_n & x_n^2 \end{pmatrix}.$$

It is worth pointing out that even if we are estimating a nonlinear function of x , we are dealing with a data observation model where x has fixed values and the variable is the parameter vector α . As a function of α , the observation model is linear.

The α that minimizes $\|\mathbf{y} - \hat{\mathbf{y}}(\mathbf{x})\|^2$ is the LS approximation

$$\hat{\alpha} = (S^\dagger S)^{-1} S^\dagger \mathbf{y}.$$

For a MATLAB example, check out the file `example.m`