

BLIND AND SEMI-BLIND MAXIMUM LIKELIHOOD METHODS FOR FIR MULTICHANNEL IDENTIFICATION

Jaouhar Ayadi, Elisabeth de Carvalho and Dirk T.M. Slock

Institut EURECOM,
B.P. 193, 06904 Sophia Antipolis Cedex, FRANCE
{ayadi, carvalho, slock}@eurecom.fr

ABSTRACT

We investigate Maximum Likelihood (ML) methods for blind and semi-blind estimation of multiple FIR channels. Two blind Deterministic ML (DML) strategies are presented. In the first one, we propose to modify the Iterative Quadratic ML (IQML) algorithm in order to "denoise" it and hence obtain consistent channel estimates. The second strategy, called Pseudo-Quadratic ML (PQML), is naturally asymptotically denoised. Links between these two approaches are established and their global convergence is proved. Furthermore, we propose semi-blind ML techniques combining PQML with two different training sequence estimation methods and compare their performance. These semi-blind techniques, exploiting the presence of known symbols, outperform their blind version. They also allow channel estimation in situations where blind and training sequence methods fail separately. Simulations are presented to demonstrate the performance of all the proposed algorithms, and comparisons between them are discussed in a blind and/or semi-blind context.

1. INTRODUCTION

Consider a sequence of symbols $a(k)$ received through m channels of length N and coefficients $h(i)$:

$$y(k) = \sum_{i=0}^{N-1} h(i)a(k \Leftrightarrow i) + v(k), \quad (1)$$

$v(k)$ is an additive independent white Gaussian noise, $rvv(k \Leftrightarrow i) = E v(k)v(i)^H = \sigma_v^2 I_m \delta_{ki}$. Assume we receive M samples, concatenated in the vector $Y_M(k)$:

$$Y_M(k) = T_M(h) A_{M+N-1}(k) + V_M(k) \quad (2)$$

$Y_M(k) = [y^H(k \Leftrightarrow M+1) \cdots y^H(k)]^H$, similarly for $V_M(k)$, and $A_M(k) = [a^H(k \Leftrightarrow M \Leftrightarrow N+2) \cdots a^H(k)]^H$, where $(\cdot)^H$ denotes hermitian transpose. The channel transfer function is $H(z) = \sum_{i=0}^{N-1} h(i)z^{-i} = [H_1^H(z) \cdots H_m^H(z)]^H$. $T_M(h)$ is a block Toeplitz matrix filled out with the channel coefficients grouped in $h = [h^H(N \Leftrightarrow 1) \cdots h^H(0)]^H$. We shall simplify the notation in (2) with $k = M \Leftrightarrow 1$ to:

$$Y = T(h) A + V. \quad (3)$$

The work of Elisabeth de Carvalho was supported by Laboratoires d'Electronique PHILIPS under contract Cifre 297/95. The work of Dirk Slock was supported by Laboratoires d'Electronique PHILIPS under contract LEP 95FAR008 and by the EURECOM Institute

We assume that $mM > M+N \Leftrightarrow 1$ in which case the channel convolution matrix $T(h)$ has more rows than columns. A channel will be said irreducible if the $H_i(z)$, $i = 1, \dots, m$ have no zeros in common, and reducible otherwise. For obvious reasons, the column space of $T(h)$ is called the signal subspace and its orthogonal complement the noise subspace.

2. BLIND DETERMINISTIC ML

The Deterministic Maximum Likelihood (DML) method was introduced for blind channel estimation in [1]. In DML, both channel coefficients and input symbols are considered as deterministic quantities, which are jointly estimated through the criterion:

$$\max_{A, h} f(Y|h) \Leftrightarrow \min_{A, h} \|Y \Leftrightarrow T(h)A\|^2 \quad (4)$$

$f(Y|h)$ is the complex probability density function (which exists as V is circular). We consider here that the blind DML identifiability conditions are verified: the channel is irreducible, the input symbols are persistently exciting and the burst sufficiently long. The channel is then identifiable up to a scale factor and we assume the regularizing constraint $\|h\| = 1$. Optimizing (4) w.r.t. A and replacing in (4), we get the following DML criterion for h :

$$\min_{\|h\|=1} Y^H P_{T(h)}^\perp Y \quad (5)$$

$P_{T(h)}^\perp$ is the orthogonal projection on the noise subspace. The key to a computationally attractive solution of the DML problem is a linear parameterization of the noise subspace.

We consider here a linear parameterization in terms of channel coefficients. Let $H^\perp(z)$ be such a parametrization; it verifies $H^\perp(z)H(z) = 0$ and $T(h^\perp)T(h) = 0$; $T(h^\perp)$ is filled with the coefficients of $H^\perp(z)$ and spans the noise subspace. An example is [2]:

$$H^\perp(z) = \begin{bmatrix} \Leftrightarrow H_2(z) & H_1(z) & 0 & \cdots & 0 \\ 0 & \Leftrightarrow H_3(z) & H_2(z) & \cdots & \vdots \\ \vdots & \vdots & \ddots & \ddots & 0 \\ H_m(z) & 0 & \cdots & 0 & \Leftrightarrow H_1(z) \end{bmatrix}. \quad (6)$$

Since $P_{T(h)}^\perp = P_{T^H(h^\perp)}$, (5) can be written as:

$$\min_{\|h\|=1} Y^H T^H(h^\perp) \mathcal{R}^+ T(h^\perp) Y \quad (7)$$

where $\mathcal{R} = T(h^\perp)T^H(h^\perp)$ and $^+$ denotes the Moore-Penrose pseudo-inverse ($T(h^\perp)$ may not be full-row rank). $T(h^\perp)$ being

linear in h , a matrix $\mathcal{Y}$ filled out with the elements of the observation vector $\mathbf{Y}$ can be found such that $\mathcal{Y}h = \mathcal{T}(h^\perp)\mathbf{Y}$. Then (8) becomes:

$$\min_{\|h\|=1} h^H \mathcal{Y}^H \mathcal{R}^+ \mathcal{Y} h \quad (8)$$

2.1. Denoised Iterative Quadratic ML (DIQML)

The Iterative Quadratic ML algorithm (IQML) solves (8) iteratively in such a way that at each step a quadratic problem appears. Indeed, at each iteration of IQML, the denominator $\mathcal{R}$, computed thanks to the previous iteration, is considered as constant and hence criterion (8) becomes quadratic. Under constraint $\|h\| = 1$, h is estimated as the minimal eigenvector of the matrix $\mathcal{Y}^H \mathcal{R}^+ \mathcal{Y}$. IQML requires a very good initialization. In [3] the channel is initialized by the SRM method [4] and the IQML estimate is proved to be consistent at high SNR. But at low SNR conditions, this method is biased because the true channel is not a stationary point of the algorithm and performs poorly even if initialized by a consistent channel estimate. We propose here a method to “denoise” the DML criterion: this denoised criterion solved in the IQML way will be consistent.

Consider the asymptotic situation where the number of data M is infinite. By the law of large numbers, the DML criterion is equivalent to its expected value which is $\text{tr}\{P_{\mathcal{T}^H(h^\perp)} E(\mathbf{Y}\mathbf{Y}^H)\}$, $E(\mathbf{Y}\mathbf{Y}^H) = \mathbf{X}\mathbf{X}^H + \sigma_v^2 \mathbf{I}$, where $\mathbf{X} = \mathcal{T}(h)A$ is the noise-free received signal. The denoising strategy consists in removing this asymptotic noise term from the DML criterion which becomes:

$$\min_{\|h\|=1} \text{tr}\{P_{\mathcal{T}^H(h^\perp)} (\mathbf{Y}\mathbf{Y}^H \Leftrightarrow \sigma_v^2 \mathbf{I})\} \quad (9)$$

Note that this operation does not change the DML criterion solution as $\sigma_v^2 \text{tr}\{P_{\mathcal{T}^H(h^\perp)}\} = \sigma_v^2 (M(m \Leftrightarrow 1) \Leftrightarrow N+1) \mathbf{I}$ is constant. (9) is solved in the IQML way: considering $\mathcal{R}$ as constant, the optimization problem is quadratic:

$$\min_{\|h\|=1} h^H \{\mathcal{Y}^H \mathcal{R}^+ \mathcal{Y} \Leftrightarrow \sigma_v^2 D\} h \quad (10)$$

where $h^H D h = \text{tr}\{\mathcal{T}^H(h^\perp) \mathcal{R}^+ \mathcal{T}(h^\perp)\}$.

Asymptotically in the number of data, DIQML is globally convergent. Indeed, asymptotically it is equivalent to the denoised criterion:

$$\min_{\|h\|=1} h^H \mathcal{X}^H \mathcal{R}^+ \mathcal{X} h \quad (11)$$

where $\mathcal{X}$ is filled out with the noise-free received signal and such that $\mathcal{X}h = \mathcal{T}(h^\perp)\mathbf{X}$. The first iteration of DIQML gives a consistent estimate: whatever the value of $\mathcal{R}$ and hence of the channel initialization, the solution of (11) is h_o , the true channel vector, (if $\text{Null}(\mathcal{R}^+) \cap \text{Im}(\mathcal{X}) = 0$, which can be easily guaranteed). A second iteration gives the global ML minimizer. Of course at high SNR global convergence is also guaranteed as it is for the original IQML.

In practice, with large but finite M , the central matrix in (10) is indefinite. So instead of removing the exact asymptotic noise term $\sigma_v^2 D$, we remove a quantity λD , which as already mentioned does not change the criterion, sufficient to render $\mathcal{Q}(h) = \mathcal{Y}^H \mathcal{R}^+ \mathcal{Y} \Leftrightarrow \lambda D$ positive semi-definite with one singularity. h is solution of:

$$\min_{\|h\|=1, \lambda} h^H \{\mathcal{Y}^H \mathcal{R}^+ \mathcal{Y} \Leftrightarrow \lambda D\} h \quad (12)$$

with constraint that $\mathcal{Q}(h)$ is positive semi-definite. h is the minimal generalized eigenvector of $\mathcal{Y}^H \mathcal{R}^+ \mathcal{Y}$ and D , and λ the generalized minimal eigenvalue. Asymptotically, $\lambda \rightarrow \sigma_v^2$, and the

criterion is equivalent to (11): asymptotic global convergence applies for h and for λ which tends to 1.

2.2. Pseudo-Quadratic ML (PQML)

The principle of PQML has been first applied to sinusoids in noise estimation [5] and then to DML in [6]. The gradient of the DML cost function may be arranged as $\mathcal{P}(h)h$, where $\mathcal{P}(h)$ is (ideally) positive semi-definite. The ML solution verifies $\mathcal{P}(h)h = 0$, which is solved under constraint $\|h\| = 1$ by the PQML strategy as follows: in a first step $\mathcal{P}(h)$ is considered constant, and as $\mathcal{P}(h)$ is positive semi-definite, the problem becomes quadratic: h is chosen in [6] as the eigenvector corresponding to the smallest absolute eigenvalue of $\mathcal{P}(h)$. This solution is used to reevaluate $\mathcal{P}(h)$ and other iterations may be done. The difficulty consists in finding the right $\mathcal{P}(h)$ and especially with the positive semi-definite constraint. In our problem:

$$\mathcal{P}(h) = \mathcal{Y}^H \mathcal{R}^+ \mathcal{Y} \Leftrightarrow \mathcal{B}^H \mathcal{B} \quad (13)$$

$\mathcal{T}^H(h^\perp)\mathcal{B} = \mathcal{B}^* h^*$ with $\mathcal{B} = [\mathcal{T}(h^\perp)\mathcal{T}^H(h^\perp)]^+ \mathcal{T}(h^\perp)\mathbf{Y}$ (* denotes the conjugate operation). Asymptotically, the effect of the second term is to remove the noise contribution present in the first one, then $\mathcal{P}(h) = \mathcal{X}^H \mathcal{R}^+ \mathcal{X}$. The asymptotic criterion is similar to the DIQML one (11) and the global convergence applies here also: again, any initialization of $\mathcal{P}(h)$ results in a consistent PQML channel estimate and the second iteration finds the global minimizer.

The matrix $\mathcal{P}(h)$ is indefinite for finite M , and applying directly the PQML strategy will not work as stated in [6], except for high SNR. PQML is closely related to IQML as the first term of (10) and (13) are the same and $E(\mathcal{B}^H \mathcal{B}) = \sigma_v^2 D$. By analogy with IQML for which $\mathcal{Q}(h)$ was also indefinite for finite M , we introduce an arbitrary λ and PQML becomes the following minimization problem:

$$\min_{\|h\|=1, \lambda} h^H \{\mathcal{Y}^H \mathcal{R}^+ \mathcal{Y} \Leftrightarrow \lambda \mathcal{B}^H \mathcal{B}\} h \quad (14)$$

with semi-definite positivity constraint on the central matrix. h is the minimal generalized eigenvector of $\mathcal{Y}^H \mathcal{R}^+ \mathcal{Y}$ and $\mathcal{B}^H \mathcal{B}$, and λ the minimal generalized eigenvalue. Asymptotically, there is global convergence for h , as described previously, and for λ .

The stationary points of PQML are the same as those of DML, this is why PQML has the same performance as DML. Asymptotically PQML gives the global ML minimizer. DIQML does not give the global ML minimizer and its performance are lower than PQML and DML.

3. SEMI-BLIND ML METHODS

Unlike purely blind approaches, semi-blind estimation techniques exploit the knowledge of certain symbols in the burst and appear superior to purely blind and training sequence methods as mentioned in [7]. Furthermore, they allow correct estimation when both blind and training sequence methods fail separately. We propose two semi-blind ML techniques both combining the previously described blind PQML and a training sequence based criterion. We will not consider DIQML, as PQML performs better. We assume that the training sequence is grouped and for simplicity reasons, is situated at the beginning of the burst. $A = [A_k^H A_u^H]^H$, where A_k groups the M_k known symbols and A_u the M_u unknown symbols of the burst.

3.1. PQML Least-Squares (PQML-LS)

This first approach remains in a deterministic perspective. We apply the DML approach to: $\mathbf{Y} = [\mathbf{Y}_{TS}^H \ \mathbf{Y}_B^H]^H$. $\mathbf{Y}_{TS} = \mathcal{T}_{TS}(h) \mathbf{A}_k + \mathbf{V}_{TS}$ groups the observations containing known symbols only. $\mathbf{Y}_B$ groups the observations containing unknown symbols: its $N \Leftrightarrow 1$ first components include a mixture of both unknown and known symbols. We do not exploit the knowledge of this symbols which will be treated as unknown. Some information is then lost. Note that this loss of information could be critical especially when the training sequence is very short, of less than N symbols!

As $\mathbf{Y}_{TS}$ and $\mathbf{Y}_B$ are decoupled in term of noise, the DML criterion for $\mathbf{Y}$ is the sum of the DML criterion for $\mathbf{Y}_{TS}$ and $\mathbf{Y}_B$:

$$\min_h \mathbf{Y}_B^H P_{\mathcal{T}^H(h)} \mathbf{Y}_B + \|\mathbf{Y}_{TS} \Leftrightarrow \mathcal{T}_{TS}(h) \mathbf{A}_k\|^2 \quad (15)$$

This semi-blind criterion is solved in the PQML way. The general semi-blind PQML strategy applies as follows: the gradient of the cost function may be written as $\mathcal{Q}(h)h + \mathcal{S}(h)$ where $\mathcal{Q}(h)$ is (ideally) positive definite. At each iteration, you suppose $\mathcal{Q}(h)$ and $\mathcal{S}(h)$ as constant, and h is the solution of a linear system, which gets used to reevaluate $\mathcal{Q}(h)$ and $\mathcal{S}(h)$ to perform other iterations. Our quantities of interest are:

$$\mathcal{Q}(h) = \mathcal{Y}_B^H \mathcal{R}_B^+ \mathcal{Y}_B \Leftrightarrow \lambda \mathcal{B}_B^H \mathcal{B}_B + \mathcal{A}_{TS}^H \mathcal{A}_{TS} \text{ and } \mathcal{S}(h) = \mathcal{A}_{TS}^H \mathbf{Y}_{TS} \quad (16)$$

where: $\mathcal{T}_{TS}(h) \mathbf{A}_k = \mathcal{A}_{TS} h$. The blind part of the criterion is solved in the PQML way and the training sequence part in the least-squares way.

We introduce the same generalized eigenvalue strategy as in the previous section which allows $\mathcal{Q}(h)$ to be positive semi-definite. At a given iteration, h has for expression:

$$h = (\mathcal{Y}_B^H \mathcal{R}_B^+ \mathcal{Y}_B \Leftrightarrow \lambda \mathcal{B}_B^H \mathcal{B}_B + \mathcal{A}_{TS}^H \mathcal{A}_{TS})^{-1} \mathcal{A}_{TS}^H \mathbf{Y}_{TS} \quad (17)$$

where the different quantities are computed thanks to the previous iteration. This criterion needs at least N known symbols to work.

3.2. PQML Weighted-Least-Squares (PQML-WLS)

The PQML-WLS mixes a deterministic and a Gaussian point of view. In the Gaussian model [7], [8], the input symbols are considered as Gaussian random variables. This hypothesis allows to robustify the problem w.r.t. to the deterministic model: an irreducible channel can be estimated up to a phase factor and for only one known symbol, not located at the edges of the burst, any channel, irreducible or not, becomes identifiable. DML can estimate reducible channels up to a scale factor, and you need at least $2N_z \Leftrightarrow 3$ known symbols, where N_z is the number of channel zeros, to identify a reducible channel [9]. Furthermore, ML based on this Gaussian model (GML) [7], [8] gives better performance than DML [9].

Let decompose $\mathbf{Y}$ as $\mathbf{Y} = [\mathbf{Y}_{SB}^H \ \mathbf{Y}_B^H]^H$. $\mathbf{Y}_B$ groups all the observations where unknown symbols only appear. $\mathbf{Y}_{SB}$ groups all the observations containing known symbols, and especially $N \Leftrightarrow 1$ observations where a mixture of both known and unknown symbols appear. The symbols in $\mathbf{Y}_B$ are treated as deterministic and DML is applied to $\mathbf{Y}_B$. In $\mathbf{Y}_{SB}$, the known symbols are treated as deterministic and the unknown symbols as i.i.d. Gaussian random variables of mean 0 and variance σ_a^2 . GML applied to $\mathbf{Y}_{SB}$ allows to take into account the known symbols that PQML-LS does not. As $\mathbf{Y}_{TS}$ and $\mathbf{Y}_B$ are decoupled in term of noise, the mixed ML criterion will be the sum of DML for $\mathbf{Y}_B$ that we solve by PQML and GML for $\mathbf{Y}_{SB}$.

$\mathbf{Y}_{SB} = \mathcal{T}_{SB}(h) \mathbf{A}_{SB} + \mathbf{V}_{SB}$, where $\mathbf{A}_{SB} = [\mathbf{A}_k^H \ \mathbf{A}_u^H]$, $\mathbf{A}_u^H$ are the unknown symbols in $\mathbf{Y}_{SB}$. $\mathbf{Y}_{SB}$ is Gaussian: $\mathbf{Y}_{SB} \sim \mathcal{N}(\mathcal{T}_{SB}(h) \mathbf{A}_{SB}^o, C_{Y_{SB} Y_{SB}})$, $\mathbf{A}_{SB}^o$ is the mean of the symbols, C_{AA} their covariance matrix, $C_{Y_{SB} Y_{SB}} = \mathcal{T}_{SB}(h) C_{AA} \mathcal{T}_{SB}^H(h) + \sigma_v^2 \mathbf{I}$. GML considers the joint estimation of h and σ_v^2 through the minimization criterion: $\max_{h, \sigma_v^2} f(\mathbf{Y} | h, \sigma_v^2)$ or:

$$\min_{h, \sigma_v^2} \{ \ln \det C_{Y_{SB} Y_{SB}} + (\mathbf{Y}_{SB} \Leftrightarrow \mathcal{T}_{SB}(h) \mathbf{A}_{SB}^o)^H C_{Y_{SB} Y_{SB}}^{-1} (\mathbf{Y}_{SB} \Leftrightarrow \mathcal{T}_{SB}(h) \mathbf{A}_{SB}^o) \} \quad (18)$$

We use PQML to solve (18). In the gradient of the cost function, the two terms coming from the derivation of $C_{Y_{SB} Y_{SB}}$ cancel each other asymptotically, and we neglect them. The quantities of interest are:

$$\mathcal{P}(h) = \mathcal{A}_{SB}^H C_{Y_{SB} Y_{SB}}^{-1} \mathcal{A}_{SB} \text{ and } \mathcal{S}(h) = C_{Y_{SB} Y_{SB}}^{-1} \mathbf{Y}_{SB} \quad (19)$$

The (approximate) PQML criterion is equivalent to the optimally WLS problem: $\min_{h, \sigma_v^2} \|\mathbf{Y} \Leftrightarrow \mathcal{T}_{SB}(h) \mathbf{A}_{SB}^o\|_{C_{Y_{SB} Y_{SB}}^{-1}}^2$. This criterion outperforms the LS criterion of the previous section: it contains all the equations of the LS criterion and allows to incorporate all the information coming from the known symbols. The mixed criterion writes:

$$\min_{h, \sigma_v^2} \mathbf{Y}_B^H P_{\mathcal{T}^H(h)} \mathbf{Y}_B + \sigma_v^2 \|\mathbf{Y} \Leftrightarrow \mathcal{T}_{SB}(h) \mathbf{A}_{SB}^o\|_{C_{Y_{SB} Y_{SB}}^{-1}}^2 \quad (20)$$

Only the information coming from the unknown symbols present in the mixed sequence previously described is lost, which is negligible as the number of observations $\mathbf{Y}_B$ will be usually large. At each iteration, the solution for h is:

$$h = (\mathcal{Y}_B^H \mathcal{R}_B^+ \mathcal{Y}_B \Leftrightarrow \lambda \mathcal{B}_B^H \mathcal{B}_B + \sigma_v^2 \mathcal{A}_{SB}^H C_{Y_{SB} Y_{SB}}^{-1} \mathcal{A}_{SB})^{-1} \mathcal{A}_{SB}^H \mathbf{Y}_{SB} \quad (21)$$

This criterion needs only one symbol to work.

PQML-WLS will give better performance than PQML-LS because the training sequence part of the criterion gives better performance. Asymptotic convergence studies are possible from two points of view. If you consider M_u as asymptotic and M_k as finite, criteria (15) and (20) are equivalent to the blind one, and inherits its convergence properties. The scale factor, not blindly identifiable, gets estimated by training sequence, this is why the estimate is not consistent. If you consider now, both M_u and M_k as infinite (with hypothesis $\frac{\sqrt{M_u}}{M_k} \rightarrow 0$), a first iteration of the algorithm gives a consistent estimate and a second one gives the global minimizer (result that can be obtained by the same asymptotic reasoning as in the blind section).

4. SIMULATION RESULTS

We consider a burst length of $M = 200$, an irreducible channel $\mathbf{H}_1$ and an ill-conditioned channel $\mathbf{H}_2$, which subchannels have one nearly common zero, of length $N = 4$ with $m = 2$ subchannels. Both channels are complex and randomly generated. The input symbols is drawn from an i.i.d. QPSK symbols sequence. The SNR is defined as $(\|\mathbf{h}\|^2 \sigma_a^2) / (m \sigma_v^2)$.

4.1. Blind algorithms

Blind estimation gives a channel estimate $\hat{\mathbf{h}}$ with $\|\hat{\mathbf{h}}\| = 1$, we adjust the right scale factor α so that $\mathbf{h}_o^H(\alpha \hat{\mathbf{h}}) = \mathbf{h}_o^H \mathbf{h}_o$ (see [7]): the

final estimate is $\hat{h} = \alpha \hat{\hat{h}}$. We plot the Normalized MSE (NMSE): $\text{NMSE} = \|\hat{h} - \hat{\hat{h}}\|^2 / \|\hat{h}\|^2$ as well as cost function of the DML criterion averaged over 100 Monte-Carlo runs in Fig.1.

We consider the channel H_1 only. The initialization of the DIQML/PQML algorithms is done by a SRM channel estimate. In Fig. 1, we illustrate the performance of DIQML and PQML at a SNR of 10 dB and 20 dB and compare it to the blind CRB computed with constraint $h_o^H \hat{h} = h_o^H h_o$ [7], corresponding to the way we have previously adjusted the scale factor. An improvement w.r.t. to the SRM initialization is clear for both algorithms and especially for PQML which outperforms DIQML. Performance can be seen to be close to the blind CRB. The evolution of the DML cost function for DIQML is not always monotonic: small fluctuations after the first iteration are observed. For the PQML algorithm the cost function was found to be always monotonic. After 1 or 2 iteration, DIQML and PQML reach their steady-state. We noticed also that the averaged least dominant generalized eigenvalue of DIQML tends to the noise variance σ_v^2 and the one of PQML to one.

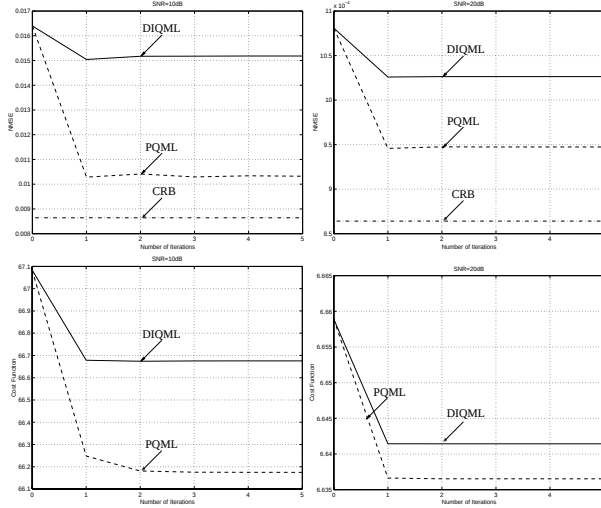


Figure 1: Performance of DIQML and PQML algorithms.

4.2. Semi-Blind algorithms

In Fig. 3 (left), we plot the averaged NMSE for PQML-LS, PQML-WLS using 10 known symbols and initialized by SRM as well as for the blind PQML for which the right scale factor is adjusted by training sequence estimation. The SNR is at 10dB and the channel is H_1 . σ_v^2 is supposed known. Both semi-blind algorithms outperforms the blind PQML and PQML-WLS outperforms PQML-LS. Differences are however small because performance for the three algorithms are already very closed to the semi-blind deterministic CRB!

In Fig. 3 (right), we present the same curves but with the ill-conditioned channel H_2 for which the blind methods fail. The algorithms are initialized by training sequence of length 10. Again performance are closed to the CRB and PQML-WLS outperforms PQML-LS. We also simulated the extreme case of channel H_2 and a training sequence of 3 for which PQML-LS cannot work at SNR=20dB. PQML-WLS converges with more difficulty than in the well-conditioned cases, but performance gets eventually closed to the semi-blind deterministic CRB. For lack of space we do not

present show the curves. We do not plot also the cost functions: they approximately reach their steady-state after one iteration, but small fluctuations could be seen afterwards.

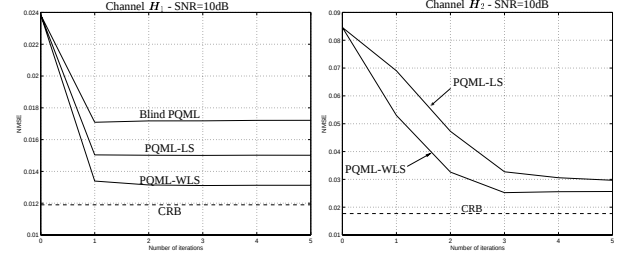


Figure 2: NMSE for PQML-LS PQML-WLS with 10 known symbols

5. REFERENCES

- [1] D.T.M. Slock. "Blind Fractionally-Spaced Equalization, Perfect-Reconstruction Filter Banks and Multichannel Linear Prediction". In *Proc. ICASSP 94 Conf.*, Adelaide, Australia, April 1994.
- [2] J. Ayadi and D. T. M. Slock. "Cramer-Rao Bounds and Methods for Knowledge Based Estimation of Multiple FIR Channels". In *Proc. SPAWC 97 Conf.*, Paris, France, April 1997.
- [3] Y. Hua. "Fast Maximum Likelihood for Blind Identification of Multiple FIR Channels". *IEEE Transactions on Signal Processing*, 44(3):661–672, March 1996.
- [4] L.A. Baccala and S. Roy. "A New Time-Domain Blind Identification Method". *IEEE Signal Processing Letters*, 1(6):89–91, June 1994.
- [5] M.R. Osborne and G.K. Smyth. "A Modified Prony Algorithm for Fitting Functions Defined by Difference Equations". *SIAM J. Sci. Stat. Comput.*, 12(2):362–382, 1991.
- [6] G. Harikumar and Y. Bresler. "Analysis and Comparative Evaluation of Techniques for Multichannel Blind Deconvolution". In *Proc. 8th IEEE Sig. Proc. Workshop Statistical Signal and Array Proc.*, pages 332–335, Corfu, Greece, June 1996.
- [7] E. de Carvalho and D. T. M. Slock. "Cramer-Rao Bounds for Semi-blind, Blind and Training Sequence based Channel Estimation". In *Proc. SPAWC 97 Conf.*, Paris, France, April 1997.
- [8] E. de Carvalho and D. T. M. Slock. "Maximum-Likelihood FIR Multi-Channel Estimation with Gaussian Prior for the Symbols". In *Proc. ICASSP 97 Conf.*, Munich, Germany, April 1997.
- [9] E. de Carvalho and D. T. M. Slock. "Asymptotic Performance of ML Methods for Semi-Blind Channel Estimation". *Proc. Asilomar Conference on Signals, Systems & Computers*, Nov. 1997.