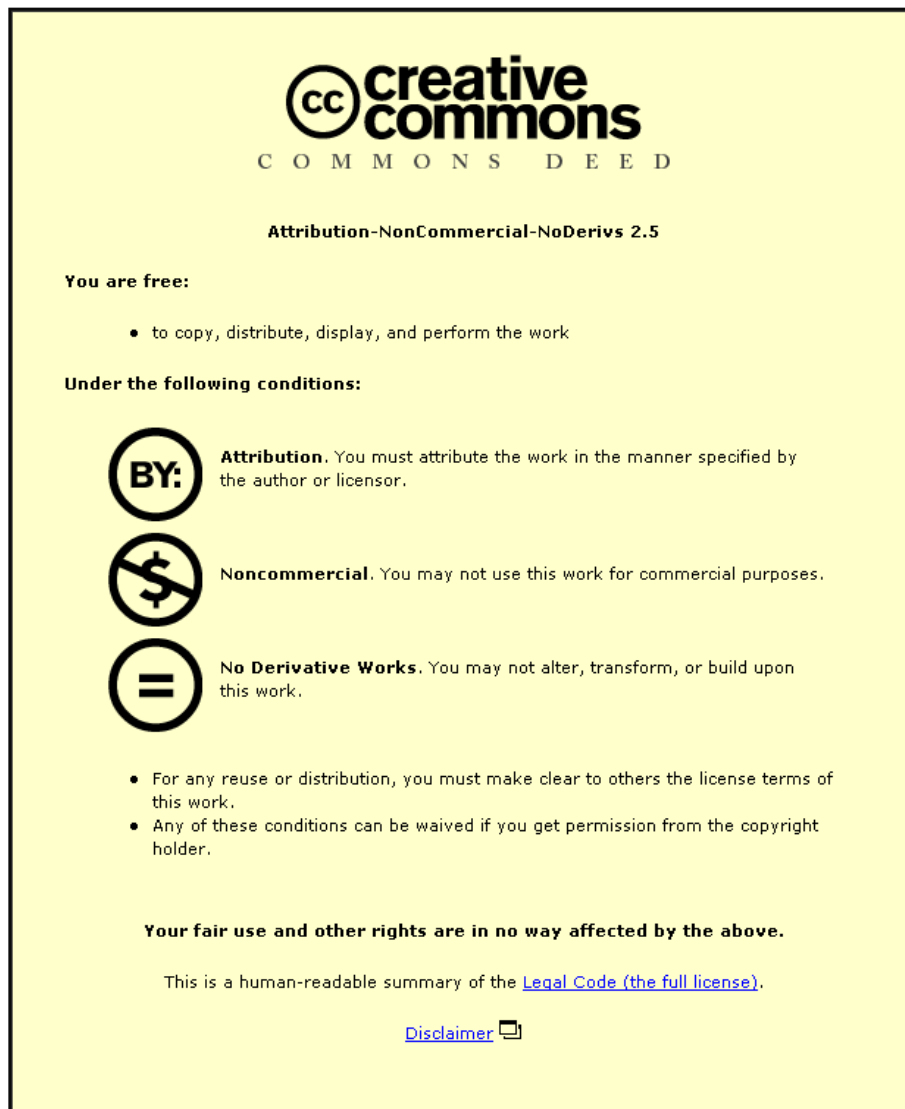


This item was submitted to Loughborough's Institutional Repository (<https://dspace.lboro.ac.uk/>) by the author and is made available under the following Creative Commons Licence conditions.



For the full text of this licence, please go to:
<http://creativecommons.org/licenses/by-nc-nd/2.5/>

VARIABLE LENGTH ADAPTIVE FILTERING WITHIN INCREMENTAL LEARNING ALGORITHMS FOR DISTRIBUTED NETWORKS

Leilei Li, Yonggang Zhang and Jonathon A. Chambers

Advanced Signal Processing Group, Department of Electronic & Electrical Engineering
Loughborough University, Loughborough LE11 3TU, UK.

(Email: l.li2@lboro.ac.uk, Y.Zhang5@lboro.ac.uk, j.a.chambers@lboro.ac.uk)

ABSTRACT

In this paper we propose the use of variable length adaptive filtering within the context of incremental learning for distributed networks. Algorithms for such incremental learning strategies must have low computational complexity and require minimal communication between nodes as compared to centralized networks. To match the dynamics of the data across the network we optimize the length of the adaptive filters used within each node by exploiting the statistics of the local signals to each node. In particular, we use a fractional tap-length solution to determine the length of the adaptive filter within each node, the coefficients of which are adapted with an incremental-learning algorithm. Simulation studies are presented to confirm the convergence properties of the scheme and these are verified by theoretical analysis of excess mean square error and mean square deviation.

Index Terms— variable tap-length, distributed processing, adaptive filters, incremental algorithm

1. INTRODUCTION

Distributed solutions, only exploiting local data exchanges and communications between immediate neighboring nodes, have been proposed [1, 2, 3] as a consequence of their reduced processing and communications requirements as compared to central solutions. The applications of such distributed adaptive networks range from sensor networks to environmental monitoring and factory instrumentation [4, 5]. However, in many applications of such distributed solutions the tap-length of the adaptive filters is assumed fixed, which is not appropriate for certain situations where the optimal tap-length is unknown or variable.

As a key parameter, tap-length plays an important role in the design of adaptive filters. It is well known that the selection of tap-length significantly influences the performance of adaptive filters: deficient tap-length is likely to result in increase of the mean-square-error (MSE); whereas the computational cost and the excess mean-square-error (EMSE) may become too high if the tap-length is too large. Since the concept of variable tap-length in adaptive filters was initially proposed in [6], various related results [7, 8, 9, 10] have been reported for the single adaptive filter case. As compared with other methods, Gong and Cowan [10] introduce a low-complexity and robust fractional tap-length (FT) algorithm based on instantaneous errors, which obtains improved convergence properties. The steady-state performance analysis of the FT algorithm are provided

in [10, 11, 12], which also provide a guideline for parameter selection in the FT algorithm.

Motivated by both the ideas of distributed estimation and variable tap-length, this paper proposes an adaptive learning algorithm which solves the parameter estimation problem in a distributed network where the tap-length of the optimal filter is not known. The steady-state performance of the new algorithm for Gaussian data is studied using weighted spatial-temporal energy conservation arguments. In particular, we derive theoretical expressions for the mean-square deviation (MSD), the EMSE and the MSE for each node within the network. Simulation studies are presented to confirm the convergence properties of the scheme and to verify the theoretical results.

The remainder of this paper is organized as follows. The proposed algorithm and its motivation are introduced in Section 2. The analysis of the fractional tap-length function is described in Section 3. Simulation results that confirm the analysis of the presented algorithm are given in section 4. Section 5 offers conclusions.

The following notations are used in this paper: boldface small and capital letters are used for random complex vectors or scalars and matrices; normal font is employed for deterministic quantities; $(\cdot)^T$ and $(\cdot)^*$ denote transposition and complex-conjugate transposition respectively; $|\cdot|^2$ and $\|\cdot\|^2$ denote the absolute squared operation and squared Euclidean norm operation.

2. ESTIMATION PROBLEM AND FORMULATION

Consider an N-node network, where we collect data and seek an unknown vector w^o , whose tap weights and tap-length are desirable to be estimated. We can decouple the adaptation rules for the tap weights and tap-length, which means that the selection of one does not depend on the other. Assuming the tap-length is L , which is estimated by the tap-length search solution that will be discussed later, each node k obtains the time observations $\{d_k(i), u_{k,i}\}$ of zero-mean complex spatial data $\{\mathbf{d}_k, \mathbf{u}_k\}$ at time instant i . Each $\mathbf{d}_k$ is a scalar value and each $\mathbf{u}_k$ is a $1 \times L$ row regression vector. In order to seek w , we formulated the linear minimum mean-square estimation problem:

$$\min_w J_L(w) \quad \text{and} \quad J_L(w) = E\|\mathbf{d} - \mathbf{U}w\|^2 \quad (1)$$

where two global matrices of the desired response and regression data are given by

$$\mathbf{d} \triangleq \text{col}\{\mathbf{d}_1, \mathbf{d}_2, \dots, \mathbf{d}_N\}, \quad (N \times 1) \quad (2)$$

$$\mathbf{U} \triangleq \text{col}\{\mathbf{u}_1, \mathbf{u}_2, \dots, \mathbf{u}_N\}, \quad (N \times L). \quad (3)$$

Let w_i be an estimate for w° at time instant i and $\psi_k(i)$ be a local estimate for w° at node k at time instant i . We introduce the incremental steepest-descent solution [2], which estimates the optimal solution w° by iterating $\psi_0(i)$ through an incremental network in the following manner:

$$\psi_k(i) = \psi_{k-1}(i) + \mu_k u_{k,i}^* (d_k(i) - u_{k,i} w_{i-1}) \quad k = 1, \dots, N \quad (4)$$

where we start from the initial condition $\psi_0(i) = w_{i-1}$ at node 1.

With the assumption of the tap-length L , we define the segmented cost function as

$$J_M(w) \triangleq E \|\mathbf{d} - \mathbf{U}_M w_M\|^2 \quad (5)$$

where $1 \leq M \leq L$, w_M and $\mathbf{U}_M$ consist of the initial M elements of w and the initial M column vectors of $\mathbf{U}$, respectively, as

$$w_M \triangleq \text{col}\{w(1), \dots, w(M)\}, \quad (1 \times M) \quad (6)$$

$$\mathbf{U}_M \triangleq \text{col}\{\mathbf{u}_1(1:M), \dots, \mathbf{u}_N(1:M)\}, \quad (N \times M). \quad (7)$$

We pose the minimum difference of mean-square errors estimation problem to seek the optimal tap-length L°

$$\min \left\{ L \mid J_{L-\Delta}(w) - J_L(w) \leq \varepsilon \right\} \quad (8)$$

where ε , predetermined by system requirements, is a small positive value and Δ is an integer value to avoid the suboptimum tap-lengths. The segmented mean-square estimation error and the mean-square estimation error at time instant i are respectively given by,

$$J_{L(i)-\Delta}(w_{i-1}) = e_{L(i)-\Delta}^2(w_{i-1}) \quad (9)$$

$$J_{L(i)}(w_{i-1}) = e_{L(i)}^2(w_{i-1}). \quad (10)$$

We start from the standard pseudo fractional tap-length implement as in [10]

$$L_f(i+1) = (L_f(i) - \alpha) + \beta \cdot [e_{L(i)-\Delta}^2(w_{i-1}) - e_{L(i)}^2(w_{i-1})] \quad (11)$$

where α and β are small positive values, α is the leakage factor used to prevent $L_f(i)$ from increasing to an undesirably large value and β is the step-size for $L_f(i)$ adaptation. Then the integer tap-length $L(i)$ is adjusted according to

$$L(i+1) = \begin{cases} \lfloor L_f(i) \rfloor & \text{if } |L(i) - L_f(i+1)| \geq \nu \\ L(i) & \text{otherwise} \end{cases} \quad (12)$$

where $\lfloor \cdot \rfloor$ rounds the embraced value to the nearest integer and the step-size parameter ν is a small integer.

Distributed networks motivate us to decompose the segment cost function and cost function as

$$J_{L-\Delta}(w) = \sum_{k=1}^N J_{k,L-\Delta}(w) \quad (13)$$

$$J_L(w) = \sum_{k=1}^N J_{k,L}(w) \quad (14)$$

where

$$J_{k,L-\Delta}(w) = E |\mathbf{d}_k - \mathbf{u}_k(1:M)w(1:M)|^2, \quad (15)$$

and

$$J_{k,L}(w) = E |\mathbf{d}_k - \mathbf{u}_k w|^2. \quad (16)$$

Such results allow us to rewrite (11) as

$$L_f(i+1) = (L_f(i) - \alpha) + \beta \cdot \sum_{k=1}^N \gamma_k(w_{i-1}) \quad (17)$$

with $\gamma_k(w_{i-1}) = e_{k,L(i)-\Delta}^2(w_{i-1}) - e_{k,L(i)}^2(w_{i-1})$. With proper choice of α and β , we will have $L(i) \rightarrow L^\circ$ as $i \rightarrow \infty$ for any initial condition, where L° is an optimal tap-length and always larger than the true tap-length of w° . Let $\ell_{k,f}(i)$ denote the local estimate of the fractional tap-length at node k at time i . α and β can be decomposed to $\alpha = \sum_{k=1}^N \alpha_k$ and $\beta = \sum_{k=1}^N \beta_k$ respectively, where α_k indicates the local leakage factor and β_k denotes the local step-size for $\ell_{k,f}(i)$ adaptation at node k . In the defined cycle, node k received the estimated fractional tap-length $\ell_{k-1,f}(i)$ from the node $k-1$. At each time instant i , we start with the initial condition $\ell_{0,f}(i) = L_f(i)$ at node 1 ($L_f(i)$ is the current global estimation for L°). At the end of the cycle, the local estimation $\ell_{N,f}(i)$ is employed as the global estimation $L_f(i+1)$ for the next time $i+1$ and the integer tap-length $L(i+1)$ is also evaluated by (12). Such implementation of a centralized solution for tap-length adaptation is described as follows:

For each time instant $i \geq 0$ repeat:
 $\ell_{0,f}(i) = L_f(i)$
 For $k=1, \dots, N$
 $\ell_{k,f}(i) = \ell_{k-1,f}(i) - \alpha_k + \beta_k \cdot \gamma_k(w_{i-1})$
 end
 $L_f(i+1) = \ell_{N,f}(i)$
 $L(i+1) = \begin{cases} \lfloor L_f(i+1) \rfloor & \text{if } |L(i) - L_f(i+1)| \geq \nu \\ L(i) & \text{otherwise} \end{cases}$
 where $\gamma_k(w_{i-1}) = e_{k,L(i)-\Delta}^2(w_{i-1}) - e_{k,L(i)}^2(w_{i-1})$

where we always hold $\ell_{k,f}(i) \geq L_{\min}$, which is the minimum tap-length. This method also requires all nodes to access the global information w_{i-1} and only adapts the integer tap-length at the end of a cycle. A fully distributed solution can be achieved by evaluating the segmented mean-square error and mean-square error from its local estimate $\psi_{k-1}^{(i)}$. This approach leads to distributed adaptation for both tap weights and tap-length. For the estimation of tap weights, a distributed version of algorithm of (4) is presented in [2] as

$$\psi_k(i) = \psi_{k-1}(i) + \mu_k u_{k,i}^* (d_k(i) - u_{k,i} \psi_{k-1}(i)) \quad k = 1, \dots, N \quad (18)$$

Let $L_k(i)$ denote the local estimate of L° at node k at time i . A distributed solution for tap-length adaptation is summarized below:

For each time instant $i \geq 0$ repeat:
 $\ell_{0,f}(i) = L_f(i)$
 For $k=1, \dots, N$
 $\ell_{k,f}(i) = \ell_{k-1,f}(i) - \alpha_k + \beta_k \cdot \gamma_k(\psi_{k-1}^{(i)})$
 $L_{k+1}(i) = \begin{cases} \lfloor \ell_{k,f}(i) \rfloor & \text{if } |L_k(i) - \ell_{k,f}(i)| \geq \nu_k \\ L_k(i) & \text{otherwise} \end{cases}$
 end
 $L_f(i+1) = \ell_{N,f}(i)$
 where $\gamma_k(\psi_{k-1}^{(i)}) = e_{k,L_k(i)-\Delta}^2(\psi_{k-1}^{(i)}) - e_{k,L_k(i)}^2(\psi_{k-1}^{(i)})$

where

$$e_{k,M(i)}(\psi_{k-1}^{(i)}) = d_k(i) - u_{k,i}(1:M(i))\psi_{k-1}^{(i)}(1:M(i)) \quad (19)$$

$$e_{k,L_k(i)}(\psi_{k-1}^{(i)}) = d_k(i) - u_{k,i}\psi_{k-1}^{(i)} \quad (20)$$

with $M(i) = L_k(i) - \Delta$. As discussed in [2], Fig. 1 illustrates that in optimization theory the distributed solution can outperform the centralized solution. The simulated curves are obtained by averaging 500 independent Monte Carlo runs with $\mu_k = 0.05$. The network utilized in the experiment has 12 nodes and seeks an unknown filter with variable tap-length $M = 10$ for $i \geq 140$ otherwise $M = 19$. The input signal is Gaussian data with $R_{u,k} = I$ and the background noise is zero mean real white Gaussian with $\sigma_{v,k}^2 = 0.001$. For both algorithms, we choose the parameters $\nu_k = \nu = 1$, $\alpha_k = 0.03$, $\beta_k = 1$, $\Delta = 4$ and $L_{\min} = L_f(0) = 6$, where the selection of parameters follows the rules in [10]-[12]. Note that, in theory, we do not need to set the upper bound for the value of fractional tap-length. However, since the length estimation uses instantaneous errors rather than averaged errors, the fractional tap-length may be, at certain time instants, at an undesired large value, which leads instantaneously to high computational and memory cost. Therefore, in practice, the upper bound is required to avoid such situation.

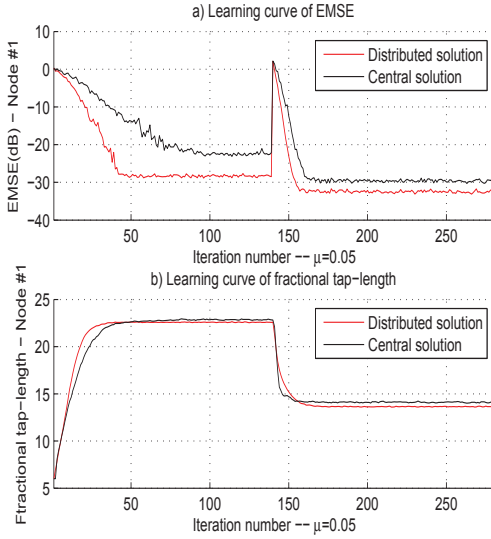


Fig. 1. Evolution curves for the distributed solution and the centralized solution at node 1: a) EMSE performance b) Fractional tap-length performance.

3. PERFORMANCE ANALYSIS

In this section, we evaluate the steady-state performance of the distributed algorithm using weighted spatial-temporal energy conservation arguments [2]. Due to space constraints, only the main procedures are given. The following assumptions are utilized

A1) In order to simplify the analysis, the estimated tap-length is assumed to be fixed at the steady-state stage.

A2) The unknown system vector w^o and $\{\mathbf{d}_k(i), \mathbf{u}_{k,i}\}$ construct:

$$\mathbf{d}_k(i) = \mathbf{u}_{k,i}w^o + \mathbf{v}_k(i) \quad (21)$$

where $\mathbf{v}_k(i)$ is a temporally and spatially white noise sequence with variance $\sigma_{v,k}^2$ and independent of $\mathbf{d}_l(j)$ for $k \neq l$ or $i \neq j$ and $\mathbf{u}_{l,j}$ for all l and j ;

A3) $\mathbf{u}_{k,i}$ is spatially and temporally independent, namely $\mathbf{u}_{k,i}$ is independent of $\mathbf{u}_{l,i}$ and $\mathbf{u}_{k,j}$ for $k \neq l$ or $i \neq j$

The following local error signals defined as in [2] are introduced to perform the evaluation:

$$\tilde{\psi}_{k-1}^{(i)} \triangleq w^o - \psi_{k-1}^{(i)}, \quad \tilde{\psi}_k^{(i)} \triangleq w^o - \psi_k^{(i)} \quad (22)$$

$$\mathbf{e}_{a,k}(i) \triangleq \mathbf{u}_{k,i}\tilde{\psi}_{k-1}^{(i)}, \quad \mathbf{e}_{p,k}(i) \triangleq \mathbf{u}_{k,i}\tilde{\psi}_k^{(i)} \quad (23)$$

Introduce further the weighted error signals:

$$\mathbf{e}_{p,k}^{\Sigma_k(i)} \triangleq \mathbf{u}_{k,i}\Sigma_k\tilde{\psi}_k^{(i)}, \quad \mathbf{e}_{a,k}^{\Sigma_k(i)} \triangleq \mathbf{u}_{k,i}\Sigma_k\tilde{\psi}_{k-1}^{(i)} \quad (24)$$

where Σ_k is a Hermitian positive-definite weighting matrix that we are free to choose at each node k . Note that the output error $\mathbf{e}_k(i) = \mathbf{e}_{a,k}(i) + \mathbf{v}_k(i)$. As a result, the steady-state quantities for each node are formed as

$$\eta_k \triangleq E\|\tilde{\psi}_{k-1}^{(\infty)}\|^2 \quad (\text{MSD}) \quad (25)$$

$$\zeta_k \triangleq E\|\tilde{\psi}_{k-1}^{(\infty)}\|_{R_{u,k}}^2 \quad (\text{EMSE}) \quad (26)$$

$$\xi_k \triangleq \zeta_k + \sigma_{v,k}^2 \quad (\text{MSE}). \quad (27)$$

where we use the weighted norm notation $\|x\|_{\Sigma}^2 = x^*\Sigma x$ for a vector x and Hermitian positive-definite $\Sigma > 0$.

As presented in [2], the spatial-temporal energy conservation relation between two successive nodes is given by

$$E\|\tilde{\psi}_k\|_{\Sigma_k}^2 = E\|\tilde{\psi}_{k-1}\|_{\Sigma'_k}^2 + \mu_k^2\sigma_{v,k}^2E\|\mathbf{u}_k\|_{\Sigma_k}^2 \quad (28)$$

where Σ'_k is given by

$$\Sigma'_k = \Sigma_k - \mu_k E(\mathbf{u}_k^*\mathbf{u}_k\Sigma_k + \Sigma_k\mathbf{u}_k^*\mathbf{u}_k) + \mu_k^2 E(\|\mathbf{u}_k\|_{\Sigma_k}^2 \mathbf{u}_k^*\mathbf{u}_k). \quad (29)$$

We assume the regressors are from a circular Gaussian distribution and introduce the eigendecomposition $R_{u,k} = Q_k\Lambda_kQ_k^*$, where Q_k is unitary and Λ_k is a diagonal matrix with the eigenvalues of $R_{u,k}$. The transformed quantities are defined as

$$\bar{\psi}_k \triangleq Q_k^*\psi_k, \quad \bar{\psi}_{k-1} \triangleq Q_k^*\psi_{k-1}, \quad \bar{\mathbf{u}}_k \triangleq \mathbf{u}_kQ_k \\ \bar{\Sigma}_k \triangleq Q_k^*\Sigma_kQ_k, \quad \bar{\Sigma}'_k \triangleq Q_k^*\Sigma'_kQ_k$$

Since Q_k is unitary, we can have $E\|\tilde{\psi}_{k-1}\|_{\Sigma_k}^2 = E\|\bar{\psi}_{k-1}\|_{\bar{\Sigma}_k}^2$ and $E\|\mathbf{u}_k\|_{\Sigma_k}^2 = E\|\bar{\mathbf{u}}_k\|_{\bar{\Sigma}_k}^2$. Let $\text{Tr}\{A\}$ denote the trace of a matrix A . Using the results for Gaussian data [13], we obtain the transformed expressions from (28) and (29)

$$E\|\bar{\psi}_k\|_{\bar{\Sigma}_k}^2 = E\|\bar{\psi}_{k-1}\|_{\bar{\Sigma}'_k}^2 + \mu_k^2\sigma_{v,k}^2\text{Tr}(\Lambda_k\bar{\Sigma}_k) \quad (30)$$

where $\bar{\Sigma}'_k$ is given by

$$\bar{\Sigma}'_k = \bar{\Sigma}_k - \mu_k X_k + \mu_k^2 Y_k \quad (31)$$

where $X_k = \Lambda_k \bar{\Sigma}_k + \bar{\Sigma}_k \Lambda_k$ and $Y_k = \Lambda_k Tr(\bar{\Sigma}_k \Lambda_k) + \tau \Lambda_k \bar{\Sigma}_k \Lambda_k$ with $\tau = 1$ for complex data and $\tau = 2$ for real data.

Therefore, we proceed as in [13, 14] by introducing the $M^2 \times 1$ vectors:

$$\delta_k = \text{vec}\{\bar{\Sigma}_k\}, \delta'_k = \text{vec}\{\bar{\Sigma}'_k\} \text{ and } \lambda_k = \text{vec}\{\Lambda_k\} \quad (32)$$

where we use the $\text{vec}\{\cdot\}$ notation in two ways: $\delta = \text{vec}\{\Sigma\}$ denotes an $M^2 \times 1$ column vector whose entries are formed by stacking the successive columns of an $M \times M$ matrix on top of each other, and $\Sigma = \text{vec}\{\delta\}$ indicates a matrix whose entries are recovered from δ . We exploit the following useful property for the $\text{vec}\{\cdot\}$ notation when working with Kronecker products: for any matrices $\{P, \Sigma, Q\}$ of compatible dimensions, it holds that

$$\text{vec}\{P\Sigma Q\} = (Q^T \otimes P)\text{vec}\{\Sigma\}. \quad (33)$$

The choice of Σ_k can make both $\bar{\Sigma}_k$ and $\bar{\Sigma}'_k$ become diagonal in (31). Applying the $\text{vec}\{\cdot\}$ operation to both sides of (31), a linear relation between the corresponding vectors $\{\delta'_k, \delta_k\}$ is obtained, namely,

$$\delta'_k = F_k \delta_k \quad (34)$$

where F_k is an $M^2 \times M^2$ matrix and given by

$$F_k = I - 2\mu_k(I \otimes \Lambda_k) + \mu_k^2(\tau(\Lambda_k \otimes \Lambda_k) + \lambda_k \lambda_k^T). \quad (35)$$

Therefore, expression (30) becomes

$$E\|\bar{\psi}_k^{(i)}\|_{\text{vec}\{\delta_k\}}^2 = E\|\bar{\psi}_{k-1}^{(i)}\|_{\text{vec}\{F_k \delta_k\}}^2 + \mu_k^2 \sigma_{v,k}^2 (\lambda_k^T \delta_k) \quad (36)$$

where we reuse the time index i for clarity. For simplicity of notation, we drop the $\text{vec}\{\cdot\}$ notation from the subscripts in (36):

$$E\|\bar{\psi}_k^{(i)}\|_{\delta_k}^2 = E\|\bar{\psi}_{k-1}^{(i)}\|_{F_k \delta_k}^2 + \mu_k^2 \sigma_{v,k}^2 (\lambda_k^T \delta_k). \quad (37)$$

Let $\rho_k = \bar{\psi}_k^{(\infty)}$, then

$$E\|\rho_k\|_{\delta_k}^2 = E\|\rho_{k-1}\|_{F_k \delta_k}^2 + \mu_k^2 \sigma_{v,k}^2 (\lambda_k^T \delta_k). \quad (38)$$

By iterating (38) over one cycle, N coupled equations are obtained:

$$E\|\rho_1\|_{\delta_1}^2 = E\|\rho_N\|_{F_1 \delta_1}^2 + g_1 \delta_1$$

$$E\|\rho_2\|_{\delta_2}^2 = E\|\rho_1\|_{F_2 \delta_2}^2 + g_2 \delta_2$$

$\vdots$

$$E\|\rho_{k-1}\|_{\delta_{k-1}}^2 = E\|\rho_{k-2}\|_{F_{k-1} \delta_{k-1}}^2 + g_{k-1} \delta_{k-1} \quad (39)$$

$$E\|\rho_k\|_{\delta_k}^2 = E\|\rho_{k-1}\|_{F_k \delta_k}^2 + g_k \delta_k \quad (40)$$

$\vdots$

$$E\|\rho_N\|_{\delta_N}^2 = E\|\rho_{N-1}\|_{F_N \delta_N}^2 + g_N \delta_N$$

with $g_k = \mu_k^2 \sigma_{v,k}^2 \lambda_k^T$. Choose the free parameters δ_k and δ_{k-1} such that $\delta_{k-1} = F_k \delta_k$ and combine (39) and (40), then we iterate this procedure across the cycle to obtain

$$\begin{aligned} E\|\rho_{k-1}\|_{\delta_{k-1}}^2 &= E\|\rho_{k-1}\|_{F_k \cdots F_N F_1 \cdots F_{k-1} \delta_{k-1}}^2 \\ &\quad + g_k F_{k+1} \cdots F_N F_1 \cdots F_{k-1} \delta_{k-1} \\ &\quad + g_{k+1} F_{k+2} \cdots F_N F_1 \cdots F_{k-1} \delta_{k-1} \\ &\quad \cdots + g_{k-2} F_{k-1} \delta_{k-1} + g_{k-1} \delta_{k-1}. \end{aligned} \quad (41)$$

Let

$$\Pi_{k-1,l} = F_{k+l-1} \cdots F_N F_1 \cdots F_{k-1}, \quad l = 1, 2, \dots, N \quad (42)$$

$$a_{k-1} = g_k \Pi_{k-1,2} + \cdots + g_{k-2} \Pi_{k-1,N} + g_{k-1} \quad (43)$$

then

$$E\|\rho_{k-1}\|_{(I - \Pi_{k-1,1})\delta_{k-1}}^2 = a_{k-1} \delta_{k-1}. \quad (44)$$

Since we are free to select the weight vector δ_{k-1} in (44), choosing $\delta_{k-1} = (I - \Pi_{k-1,1})^{-1} q$ or $\delta_{k-1} = (I - \Pi_{k-1,1})^{-1} \lambda_k$ results in the expressions for the steady-state MSD, EMSE and MSE at node k :

$$\eta_k = E\|\rho_{k-1}\|_q^2 = a_{k-1} (I - \Pi_{k-1,1})^{-1} q \quad (\text{MSD}) \quad (45)$$

$$\zeta_k = E\|\rho_{k-1}\|_{\lambda_k}^2 = a_{k-1} (I - \Pi_{k-1,1})^{-1} \lambda_k \quad (\text{EMSE}) \quad (46)$$

$$\xi_k = \zeta_k + \sigma_{v,k}^2 \quad (\text{MSE}) \quad (47)$$

where $q = \text{vec}\{I\}$ and $\lambda_k = \text{vec}\{\Lambda_k\}$.

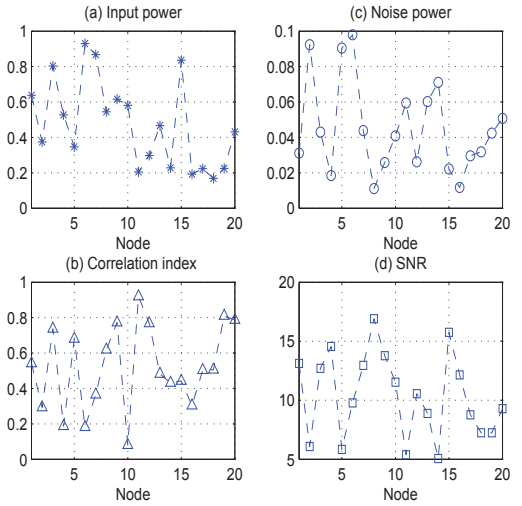


Fig. 2. Node profile: a) Input power b) Noise power c) Correlation index d) SNR.

4. SIMULATIONS

In this section, we compare the theoretical performance with the computer simulations in a system identification scenario. All simulations results are averaged over 100 independent Monte Carlo runs. The steady-state curves are obtained by averaging the last 2000 instantaneous samples of 20,000 iterations. We considered a network with 20 nodes seeking an unknown filter with $M = 10$ taps. A correlated Gaussian signal is used to generate the inputs at each node k which satisfies the recursion

$$u_k(i) = a_k u_k(i-1) + b_k \cdot c_k(i). \quad (48)$$

Expression (48) produces a first-order autoregressive (AR) process with a pole at a_k ; c_k is a white, zero-mean, Gaussian random sequence with unity variance, $a_k \in (0, 1]$ and $b_k = \sqrt{\sigma_{u,k}^2 \cdot (1 - a_k^2)}$.

In this way, the covariance matrix $R_{u,k}$ of the regressor $\mathbf{u}_{k,i}$ is a 10×10 Toeplitz matrix with entries $r_k(m) = \sigma_{u,k}^2 a_k^{|m|}$, $m = 0, \dots, M-1$ with $\sigma_{u,k}^2 \in [0, 1)$. The background noise has variance $\sigma_{v,k}^2 \in [0, 0.1)$ across the network. The statistical profiles are illustrated in Fig. 2. Fig. 3 and Fig. 4 show that the theoretical results for MSE match well with the simulated results, where $\nu_k = 1$, $\alpha_k = 0.05$, $\beta_k = 1$, $\Delta = 3$, and $L_{\min} = L_f(0) = 4$. Similar results for EMSE and MSD have been obtained but are not included due to space limitations.

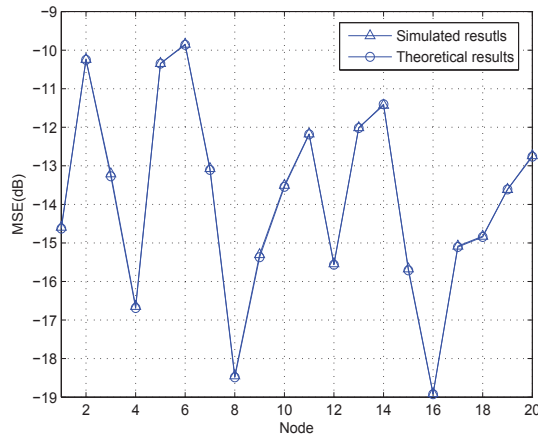


Fig. 3. MSE versus node $k - \mu_k = 0.02$.

5. CONCLUSION

In this paper, a new algorithm has been proposed for structure adaptation of adaptive filters in an incremental distributed network. Under the assumptions A1, A2 and A3, weighted spatial-temporal energy conservation argument are used to analyze steady-state mean square performance in the Gaussian case. Numerical simulations show that there is a good match between simulated results and theoretical results. Future work includes more sophisticated distributed estimation in a cooperative fashion.

6. REFERENCES

- [1] A. H. Sayed and C. G. Lopes, "Distributed recursive least-squares strategies over adaptive networks," *Proc. 40th Asilomar Conference on Signals, Systems and Computers*, pp. 223–237, Oct. 2006.
- [2] C. G. Lopes and A. H. Sayed, "Incremental adaptive strategies over distributed networks," *IEEE Trans. Signal Processing*, vol. 55, no. 8, pp. 4064–4077, 2007.
- [3] L. Li and J. A. Chambers, "A new incremental affine projection based adaptive algorithm for distributed networks," *Fast Communication in Signal Processing*, vol. 88, no. 10, pp. 2599–2603, 2008.
- [4] D. Estrin, L. Girod, G. Pottie, and M. Srivastava, "Instrumenting the world with wireless sensor networks," *Proc. ICASSP*, pp. 2033–2036, May. 2001.
- [5] M. G. Rabbat and R. D. Nowak, "Quantized incremental algorithms for distributed optimization," *IEEE Journal on Selected Areas in Communications*, vol. 23, no. 4, pp. 798–808, Apr. 2007.
- [6] Z. Pritzker and A. Feuer, "Variable length stochastic gradient algorithm," *IEEE Trans. Signal Process.*, vol. 39, no. 4, pp. 977–1001, Apr. 1991.
- [7] Y. K. Won, R. H. Park, J. H. Park, and B. U. Lee, "Variable LMS algorithm using the time constant concept," *IEEE Trans. Consumer Electron.*, vol. 40, no. 3, pp. 655–661, 1994.
- [8] F. Riera-Palou, J. M. Norsa, and D. J. M. Cruickshank, "Linear equalizers with dynamic and automatic length selection," *Electron. Lett.*, vol. 37, no. 25, pp. 1553–1554, Dec. 2001.
- [9] Y. Gu, K. Tang, H. Cui, and W. Du, "LMS algorithm with gradient descent filter length," *IEEE Signal Process. Lett.*, vol. 11, no. 3, pp. 305–307, Mar. 2004.
- [10] Y. Gong and C. F. N. Cowan, "An LMS style variable tap-length algorithm for structure adaptation," *IEEE Tran. Signal Processing*, vol. 53, no. 7, pp. 2400–2407, 2005.
- [11] Y. Zhang, N. Li, J. Chambers, and A. H. Sayed, "Steady state performance analysis of a variable tap-length LMS algorithm," *IEEE Transactions on Signal Processing*, vol. 56, no. 2, pp. 839–845, 2008.
- [12] L. Li and J. A. Chambers, "A novel adaptive leakage factor scheme for enhancement of a variable tap-length learning algorithm," *ICASSP 2008*, pp. 2837–2840, Apr. 2008.
- [13] A. H. Sayed, *Fundamentals of Adaptive Filtering*, John Wiley & Sons, Inc., Hoboken, NJ, 2003.
- [14] L. Li, Y. Zhang, and J. A. Chambers, "Steady-state performance of incremental learning over distributed networks for non-Gaussian data," to appear in *ICSP 2008*, Oct. 2008.

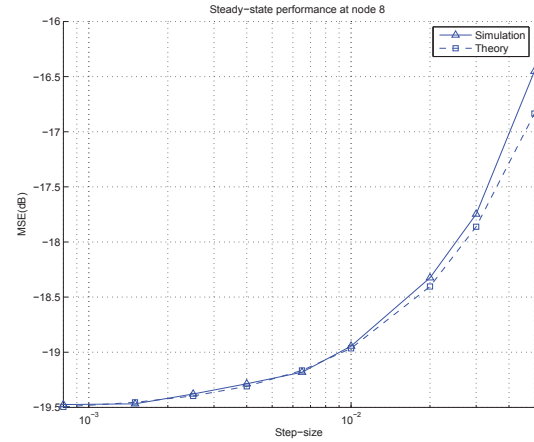


Fig. 4. MSE versus step-size at node 8.