

Distributed Bayesian vector estimation using noise-optimized low-resolution sensor observations

Jia Liu^a, Fabing Duan^{a,*}, François Chapeau-Blondeau^b, Derek Abbott^c

^a Institute of Complexity Science, Qingdao University, Qingdao 266071, PR China

^b Laboratoire Angevin de Recherche en Ingénierie des Systèmes (LARIS), Université d'Angers, 62 avenue Notre Dame du Lac, 49000 Angers, France

^c Centre for Biomedical Engineering (CBME) and School of Electrical & Electronic Engineering, The University of Adelaide, Adelaide, SA 5005, Australia

ARTICLE INFO

Article history:

Available online 6 September 2021

Keywords:

Distributed Bayesian estimation

Low-resolution sensor

Optimal added noise

Noise-enhanced vector estimator

Stochastic resonance

ABSTRACT

The distributed Bayesian vector parameter estimation problem based on low-resolution observations is investigated in a network, where each node represents an ensemble of estimates from a large number of sensors. A noise-enhanced Bayesian vector estimator that benefits from artificially added noise is proposed. For a network node composed of a sufficiently large number of identical low-resolution sensors, a lemma governing the weight coefficients is proven, and low-cost calculation expressions of the designed estimator and its Bayesian mean square error (MSE) are derived by avoiding costly computations due to high-dimensional matrix inversions. Experimental results show that by intentionally adding an appropriate amount of noise to networks of the low-resolution sensors, the MSE of the designed Bayesian vector estimator can be significantly reduced.

© 2021 Elsevier Inc. All rights reserved.

1. Introduction

The widespread use of low-power and low-complexity sensors in wireless networks and digital communication systems offers an answer to the demanding requirements of cost constraints and bandwidth limitations [1–4], but poses the challenge of achieving efficient processing from low-resolution observations of many sensors [5–7]. In recent years, increasing interest has focused on studying efficient parameter estimation based on low-resolution observations. In order to reduce the signal distortion and improve the accuracy of parameter estimation, the operation of adding optimal noise before sensors is often employed in practical applications such as audio coding [8], image compression [9], distributed estimation [2–4,10,11], direction-of-arrival [12], and multiple-input multiple-output communications [13], where it is also commonly referred to as dithering [1,5,7,14]. Dithering has been widely investigated for reducing the mean square error (MSE) of deterministic or random parameter estimation [7,10,13,15], or for minimizing the Cramér-Rao lower bound (CRLB) [3,5,6], by means of an optimal amount of added noise. Stochastic resonance [16,17] also represents a broader class of noise-aided phenomena, and has been exploited for noise-improved information processing for var-

ious operations, including signal transmission [18–20], detection [21–24] and estimation [19,25,26].

In the field of signal estimation, Zeitler et al. [5] proved that, for a single-bit dithered quantizer and a Gaussian prior, any estimator using the quantizer output sequence is asymptotically at least $10 \log(\pi/2) \approx 1.97$ dB worse than a minimum MSE (MMSE) estimator using unquantized observations. From quantized observations in multiplicative noise environments, the enhancements of the estimation accuracy by multiplicative noise were reported in [27,28]. For the linear MMSE (LMMSE) case, the optimal added noise for the initial rate of the noise benefit in binary quantizer arrays [29] and for minimizing the MSE of the nonlinear transformation [30] were derived. Using quantized observations of an M-level quantizer, the optimal added noise is proven to be a constant signal level in terms of minimizing the Bayesian CRLB [31]. Recently, an interesting approach to adding noise for signal estimation has been proposed employing an ensemble of estimates from a sufficiently large number of sensors [32–34], whose inputs contain the same signal perturbed by mutually independent noise components but with low-resolution outputs. This ensemble integration method of sensor outputs is asymptotically equivalent to the expectation of estimates with respect to the intentionally added noise distribution [32]. Based on this deduction, these noise-enhanced estimators [32,34] demonstrated effective improvement in the estimation accuracy of a random parameter and in experiments with trained feedforward neural networks [33,35] thus il-

* Corresponding author.

E-mail address: fabing.duan@gmail.com (F. Duan).

Table 1
Summary of notation.

Symbol	Notation	Symbol	Notation
x	scalar measurement	$\mathbf{x}$	observation vector
θ	random parameter	$\boldsymbol{\theta}$	parameter vector
$E_{\theta}(\boldsymbol{\theta})$	parameter mean vector	$\mathbf{C}_{\theta\theta}$	covariance matrix of $\boldsymbol{\theta}$
$\mathbf{h}$	row vector of $\mathbf{H}$	$\mathbf{H}$	observation matrix
ξ	background noise	ξ	background noise vector
g	sensor transfer function	w_0	bias weight coefficient
η	added noise	$\boldsymbol{\eta}$	added noise vector
y	sensor output	$\mathbf{y}$	node output vector
ω	weight coefficient for each sensor	$\boldsymbol{\omega}$	network weight vector
$\mathbf{w}$	weight coefficient for each node	$\mathbf{w}$	network weight vector
$\mathbf{z}$	network output vector	f_{η}	PDF of added noise η
$\mathbf{C}_{\theta\mathbf{z}}$	correlation vector between $\boldsymbol{\theta}$ and $\mathbf{z}$	$\mathbf{C}_{\mathbf{z}\mathbf{z}}$	covariance matrix of $\mathbf{z}$
$\hat{\theta}_i$	Bayesian estimator	$B_{\hat{\theta}_i}$	Bayesian MSE of $\hat{\theta}_i$
$\bar{y}$	average output of node	$\bar{\mathbf{y}}$	network output vector
M	sensor number	N	network node number
g^*	convergence of $\bar{y}$ for $M \rightarrow \infty$	$\mathbf{g}^*$	network output vector for $M \rightarrow \infty$
$\hat{\boldsymbol{\theta}}^M$	Bayesian vector estimator with M identical sensors	$\mathbf{B}_{\text{NE}}^M$	Bayesian MSE matrix of $\hat{\boldsymbol{\theta}}^M$
$\hat{\boldsymbol{\theta}}^*$	vector estimator of $\boldsymbol{\theta}$ with infinite identical sensors	$\mathbf{B}_{\text{NE}}^*$	Bayesian MSE matrix of $\hat{\boldsymbol{\theta}}^*$
$\hat{\boldsymbol{\theta}}_{\text{LMMSE}}$	LMMSE estimator on the measurement $\mathbf{x}$	$\mathbf{B}_{\text{LMMSE}}$	Bayesian MSE matrix of $\hat{\boldsymbol{\theta}}_{\text{LMMSE}}$
$E_{\mathbf{x}}(\cdot)$	expectation with respect to the PDF $f_{\mathbf{x}}(\mathbf{x})$	$E_{\mathbf{x},\mathbf{y}}(\cdot)$	expectation with respect to the joint PDF $f_{\mathbf{x},\mathbf{y}}(\mathbf{x}, \mathbf{y})$
$\mathbf{1}_M$	$M \times 1$ -dimensional unit vector	σ^2	variance of the random variable
$\mathbf{I}_N$	$N \times N$ -dimensional unit matrix	γ	quantizer threshold

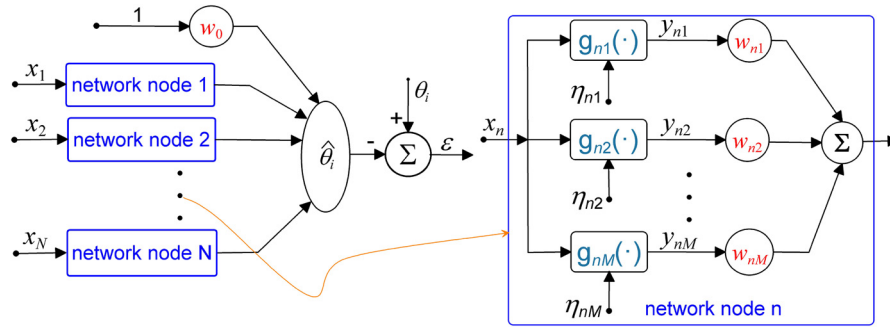


Fig. 1. Block diagram representation of the noise-enhanced estimator $\hat{\theta}_i$ ($i = 1, 2, \dots, p$) with N network nodes, and each node composed of M sensors and M mutually independent added noise components η_{nm} for $n = 1, 2, \dots, N$ and $m = 1, 2, \dots, M$.

illustrating the benefit of an optimal amount of added noise to the function approximation.

Most of studies of noise-enhanced parameter estimation focus on estimating a scalar parameter [5,27–30,32,34]. In many problems of interest we need to estimate a random vector of parameters, e.g., estimating a flow vector from the measurement of hydraulic fields [36–38], the impulse response of a linear system model [39] and target tracking [40]. Then, the noise-enhanced vector parameter or multiple-parameter estimator also triggered various studies [31,41,42] in an M-level or binary quantizer. Another practical vector parameter estimation problem arising from sensor network applications is based on a set of spatially distributed sensors, i.e. distributed parameter estimation [2–4]. For this estimation problem, there are two meaningful questions worthy to be answered: how to exploit the artificially added noise to design a vector estimator based on the spatially distributed low-resolution sensors, and to what extent the Bayesian MSE of the designed estimator can be improved by the optimal added noise.

In this paper, we design a Bayesian vector parameter estimator, as represented in Fig. 1, to exploit mutually independent added noise components in the low-resolution network for improving estimation accuracy effectively. We first prove that, for any two identical sensors in a network node, the optimum weight coefficients assigned to two sensors are equal to each other. Then, when each network node is composed of a sufficiently large number of identical sensors, this property leads to significantly simplified representations of the estimator and its Bayesian MSE matrix by

avoiding the costly computation of high-dimensional matrix inversions. For a prior Gaussian joint probability distribution of the random parameters and background noise, the optimal probability density function (PDF) of added noise that minimizes the trace of the Bayesian MSE matrix is theoretically characterized by an approximate expression. When the joint probability distribution is non-Gaussian, the improvement of the designed estimator by the added Gaussian noise is numerically evaluated by simulated realizations of the time averaging algorithm. Theoretical and numerical analyses show that, for an optimal added noise level or optimal added noise type, the Bayesian MSE of the designed estimator is very close to that obtained by the MMSE (or LMMSE) Bayesian estimator on the original measurements over a wide range of the input signal-to-noise ratio (SNR). The main contributions of this paper include: proposing a distributed Bayesian vector estimator based on a low-complexity sensor network and confirming the extent to which intentionally artificially added noise reduces the Bayesian MSE of the proposed estimator. List of symbols in this paper is shown in Table 1.

2. Model and formulation

Consider a sensor network, as shown in Fig. 1, where each network node receives a scalar measurement sequence

$$x_n = \mathbf{h}_n \boldsymbol{\theta} + \xi_n, \quad n = 1, 2, \dots, N. \quad (1)$$

Here, $\mathbf{h}_n$ is the $1 \times p$ row vector of the $N \times p$ ($N > p$) known deterministic observation matrix $\mathbf{H} = [\mathbf{h}_1^\top, \mathbf{h}_2^\top, \dots, \mathbf{h}_N^\top]^\top$, $\boldsymbol{\theta} = [\theta_1, \theta_2, \dots, \theta_p]^\top$ is a $p \times 1$ unknown random vector parameter with mean $\mathbf{E}_\theta(\boldsymbol{\theta})$ and covariance matrix $\mathbf{C}_{\theta\theta}$, and the zero-mean background noise variable ξ_n is uncorrelated with $\boldsymbol{\theta}$. Let $\mathbf{x} = [x_1, x_2, \dots, x_N]^\top$ and $\boldsymbol{\xi} = [\xi_1, \xi_2, \dots, \xi_N]^\top$, we get the linear vector measurement model $\mathbf{x} = \mathbf{H}\boldsymbol{\theta} + \boldsymbol{\xi}$ [39,40].

As illustrated in Fig. 1, the n -th network node has a common input x_n , but consists of an ensemble of M sensors with low-resolution outputs

$$y_{nm} = g_{nm}(x_n + \eta_{nm}), \quad m = 1, 2, \dots, M, \quad (2)$$

where g_{nm} is the transfer function (e.g. quantizer in Eq. (21) or hard-limiter in Eq. (22)) of the m -th sensor in the node n , and M mutually independent added noise components η_{nm} are with a common PDF f_η , which is assumed to be designable. Then, we consider the problem of estimating the unknown random vector parameter $\boldsymbol{\theta}$ from sensor observation vectors $\mathbf{y}_n = [y_{n1}, y_{n2}, \dots, y_{nM}]^\top$ for $n = 1, 2, \dots, N$.

First, as indicated in Fig. 1, we assign an adjustable weight ω_{nm} to each sensor output y_{nm} and denote the weight vector as $\boldsymbol{\omega}_n = [\omega_{n1}, \omega_{n2}, \dots, \omega_{nM}]^\top$, and then the n -th network output is given by $\boldsymbol{\omega}_n^\top \mathbf{y}_n$. An intermediate linear estimator is designed as

$$\hat{\theta}_i = w_0 + \boldsymbol{\omega}^\top \mathbf{z} \quad (3)$$

to estimate the i -th parameter θ_i ($i = 1, 2, \dots, p$), where w_0 is the biasing weight, the $MN \times 1$ weight vector $\boldsymbol{\omega} = [\omega_1^\top, \omega_2^\top, \dots, \omega_N^\top]^\top$ and the $MN \times 1$ network output vector $\mathbf{z} = [\mathbf{y}_1^\top, \mathbf{y}_2^\top, \dots, \mathbf{y}_N^\top]^\top$. Then, the error $\varepsilon_i = \theta_i - \hat{\theta}_i$ and the Bayesian MSE of the intermediate estimate $\hat{\theta}_i$ becomes

$$\mathbf{E}_{\mathbf{x}, \eta}(\varepsilon_i^2) = \mathbf{E}_{\mathbf{x}, \eta}[(\theta_i - w_0 - \boldsymbol{\omega}^\top \mathbf{z})^2]. \quad (4)$$

Setting the derivative $\partial \mathbf{E}_{\mathbf{x}, \eta}(\varepsilon_i^2) / \partial w_0 = 0$, we find the optimum biasing weight as $w_0^{\text{opt}} = \mathbf{E}_{\theta_i}(\theta_i) - \boldsymbol{\omega}^\top \mathbf{E}_{\mathbf{x}, \eta}(\mathbf{z})$. Then, Eq. (3) can be rewritten as $\hat{\theta}_i = \mathbf{E}_{\theta_i}(\theta_i) + \boldsymbol{\omega}^\top [\mathbf{z} - \mathbf{E}_{\mathbf{x}, \eta}(\mathbf{z})]$ and the Bayesian MSE of $\hat{\theta}_i$ in Eq. (4) is obtained as

$$\begin{aligned} \mathbf{E}_{\mathbf{x}, \eta}(\varepsilon_i^2) &= \mathbf{E}_{\mathbf{x}, \eta} \left[\left(\theta_i - \mathbf{E}_{\theta_i}(\theta_i) - \boldsymbol{\omega}^\top [\mathbf{z} - \mathbf{E}_{\mathbf{x}, \eta}(\mathbf{z})] \right)^2 \right] \\ &= \mathbf{C}_{\theta_i} - 2\boldsymbol{\omega}^\top \mathbf{C}_{\theta_i \mathbf{z}} + \boldsymbol{\omega}^\top \mathbf{C}_{\mathbf{z} \mathbf{z}} \boldsymbol{\omega}, \end{aligned} \quad (5)$$

where the variance of the parameter θ_i is $\mathbf{C}_{\theta_i} = \mathbf{E}_{\theta_i}[(\theta_i - \mathbf{E}_{\theta_i}(\theta_i))^2]$, the covariance matrix $\mathbf{C}_{\mathbf{z} \mathbf{z}} = \mathbf{E}_{\mathbf{x}, \eta}[(\mathbf{z} - \mathbf{E}_{\mathbf{x}, \eta}(\mathbf{z}))(\mathbf{z} - \mathbf{E}_{\mathbf{x}, \eta}(\mathbf{z}))^\top]$ and the correlation vector $\mathbf{C}_{\theta_i \mathbf{z}} = \mathbf{E}_{\mathbf{x}, \eta}[(\mathbf{z} - \mathbf{E}_{\mathbf{x}, \eta}(\mathbf{z}))(\theta_i - \mathbf{E}_{\theta_i}(\theta_i))]$. It is clear from this expression that the Bayesian MSE $\mathbf{E}_{\mathbf{x}, \eta}(\varepsilon_i^2)$ is precisely a quadratic function of the weight vector $\boldsymbol{\omega}$ when the vector parameter $\boldsymbol{\theta}$, the background vector noise $\boldsymbol{\xi}$ and sensor noise components η_{nm} are stationary stochastic variables. Thus, to obtain the minimum $\mathbf{E}_{\mathbf{x}, \eta}(\varepsilon_i^2)$ with respect to $\boldsymbol{\omega}$, we set $\partial \mathbf{E}_{\mathbf{x}, \eta}(\varepsilon_i^2) / \partial \boldsymbol{\omega} = -2\mathbf{C}_{\theta_i \mathbf{z}} + 2\mathbf{C}_{\mathbf{z} \mathbf{z}} \boldsymbol{\omega} = 0$ and find the optimum weight vector, i.e. Wiener weight vector $\boldsymbol{\omega}^{\text{opt}} = \mathbf{C}_{\mathbf{z} \mathbf{z}}^{-1} \mathbf{C}_{\theta_i \mathbf{z}}$. Substituting $\boldsymbol{\omega}^{\text{opt}}$ into Eq. (3) and noting $\mathbf{C}_{\theta_i \mathbf{z}} = \mathbf{C}_{\mathbf{z} \theta_i}^\top$, we can re-express the estimator $\hat{\theta}_i$ in Eq. (3) as

$$\hat{\theta}_i = \mathbf{E}_{\theta_i}(\theta_i) + \mathbf{C}_{\theta_i \mathbf{z}} \mathbf{C}_{\mathbf{z} \mathbf{z}}^{-1} [\mathbf{z} - \mathbf{E}_{\mathbf{x}, \eta}(\mathbf{z})]. \quad (6)$$

The corresponding Bayesian MSE of $\hat{\theta}_i$ is given by $B_{\hat{\theta}_i} = \mathbf{E}_{\mathbf{x}, \eta}[(\theta_i - \hat{\theta}_i)^2] = \mathbf{C}_{\theta_i} - \mathbf{C}_{\theta_i \mathbf{z}} \mathbf{C}_{\mathbf{z} \mathbf{z}}^{-1} \mathbf{C}_{\mathbf{z} \theta_i}$. It is seen that, when the sensor number M is large, the performance evaluation of the estimator $\hat{\theta}_i$ in Eq. (6) requires the costly computations of the $MN \times 1$ correlation vector $\mathbf{C}_{\theta_i \mathbf{z}}$ and the $MN \times MN$ inverse matrix $\mathbf{C}_{\mathbf{z} \mathbf{z}}^{-1}$. In order to reduce the calculation cost, we will show that, for a number of identical sensors in each network node, simplifications of the designed estimator $\hat{\theta}_i$ and its performance can be devised.

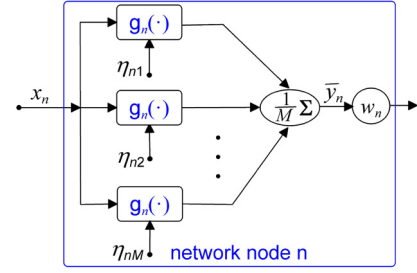


Fig. 2. Block diagram representation of the n -th network nodes with M identical sensors $g_{nm}(\cdot) = g_n(\cdot)$ and M mutually independent added noise components η_{nm} for $m = 1, 2, \dots, M$. Only one weight coefficient w_n is assigned to the average node output $\bar{y}_n$.

Theorem 1. For M identical transfer functions $g_{nm}(\cdot) = g_n(\cdot)$ of sensors in the n -th network node, the $M \times 1$ optimum weight vector $\boldsymbol{\omega}_n^{\text{opt}}$ of the network node has M equivalent weight coefficients.

Proof of Theorem 1 is presented in Appendix A. Based on Theorem 1, when the transfer functions $g_{nm}(\cdot)$ of sensors are identical in the n -th network node, M weight coefficients w_{nm} are reduced to only one weight coefficient w_n^{opt} indicated in Eq. (A.10). Thus, the estimator $\hat{\theta}_i$ in Eq. (6) and its performance evaluation of $B_{\hat{\theta}_i}$ can be calculated with low computational complexities as follows.

Corollary 1. When the transfer functions $g_{nm}(\cdot) = g_n(\cdot)$ of sensors are identical in the n -th network node, each network node model can be simplified as Fig. 2 with only one weight coefficient w_n assigned to the average node output

$$\bar{y}_n = \frac{1}{M} \sum_{m=1}^M g_n(x_n + \eta_{nm}) \quad (7)$$

for $n = 1, 2, \dots, N$. Letting the weight vector $\mathbf{w} = [w_1, w_2, \dots, w_N]^\top$ and the network output vector $\bar{\mathbf{y}} = [\bar{y}_1, \bar{y}_2, \dots, \bar{y}_N]^\top$, the designed estimator $\hat{\theta}_i$ in Eq. (6) can be simplified as

$$\hat{\theta}_i^M = \mathbf{E}_{\theta_i}(\theta_i) + \mathbf{C}_{\theta_i \bar{\mathbf{y}}} \mathbf{C}_{\bar{\mathbf{y}} \bar{\mathbf{y}}}^{-1} [\bar{\mathbf{y}} - \mathbf{E}_{\mathbf{x}, \eta}(\bar{\mathbf{y}})] \quad (8)$$

with its Bayesian MSE

$$B_{\hat{\theta}_i^M} = \mathbf{E}_{\mathbf{x}, \eta}[(\theta_i - \hat{\theta}_i^M)^2] = \mathbf{C}_{\theta_i} - \mathbf{C}_{\theta_i \bar{\mathbf{y}}} \mathbf{C}_{\bar{\mathbf{y}} \bar{\mathbf{y}}}^{-1} \mathbf{C}_{\bar{\mathbf{y}} \theta_i}, \quad (9)$$

where the $1 \times N$ correlation vector $\mathbf{C}_{\theta_i \bar{\mathbf{y}}} = \mathbf{E}_{\mathbf{x}, \eta}[(\theta_i - \mathbf{E}_{\theta_i}(\theta_i))(\bar{\mathbf{y}} - \mathbf{E}_{\mathbf{x}, \eta}(\bar{\mathbf{y}}))^\top]$ and the $N \times N$ covariance matrix $\mathbf{C}_{\bar{\mathbf{y}} \bar{\mathbf{y}}} = \mathbf{E}_{\mathbf{x}, \eta}[(\bar{\mathbf{y}} - \mathbf{E}_{\mathbf{x}, \eta}(\bar{\mathbf{y}}))(\bar{\mathbf{y}} - \mathbf{E}_{\mathbf{x}, \eta}(\bar{\mathbf{y}}))^\top]$. The p designed estimator $\hat{\theta}_i^M$ in Eq. (8) can be combined into a vector estimator as

$$\hat{\boldsymbol{\theta}}_{\text{NE}}^M = [\hat{\theta}_1^M, \hat{\theta}_2^M, \dots, \hat{\theta}_p^M]^\top = \mathbf{E}_\theta(\boldsymbol{\theta}) + \mathbf{C}_{\boldsymbol{\theta} \bar{\mathbf{y}}} \mathbf{C}_{\bar{\mathbf{y}} \bar{\mathbf{y}}}^{-1} [\bar{\mathbf{y}} - \mathbf{E}_{\mathbf{x}, \eta}(\bar{\mathbf{y}})], \quad (10)$$

where the $p \times N$ correlation matrix $\mathbf{C}_{\boldsymbol{\theta} \bar{\mathbf{y}}} = \mathbf{C}_{\bar{\mathbf{y}} \boldsymbol{\theta}}^\top = [\mathbf{C}_{\theta_1 \bar{\mathbf{y}}}, \mathbf{C}_{\theta_2 \bar{\mathbf{y}}}, \dots, \mathbf{C}_{\theta_p \bar{\mathbf{y}}}]^\top$. The Bayesian matrix of the vector estimator $\hat{\boldsymbol{\theta}}_{\text{NE}}^M$ is given by

$$\mathbf{B}_{\text{NE}}^M = \mathbf{E}_{\mathbf{x}, \eta}[(\boldsymbol{\theta} - \hat{\boldsymbol{\theta}}_{\text{NE}}^M)(\boldsymbol{\theta} - \hat{\boldsymbol{\theta}}_{\text{NE}}^M)^\top] = \mathbf{C}_{\boldsymbol{\theta} \boldsymbol{\theta}} - \mathbf{C}_{\boldsymbol{\theta} \bar{\mathbf{y}}} \mathbf{C}_{\bar{\mathbf{y}} \bar{\mathbf{y}}}^{-1} \mathbf{C}_{\bar{\mathbf{y}} \boldsymbol{\theta}}, \quad (11)$$

whose diagonal element $[\mathbf{B}_{\text{NE}}^M]_{ii} = B_{\hat{\theta}_i^M}$ for $i = 1, 2, \dots, p$.

Proof of Corollary 1 is presented in Appendix B. It is seen that the computation cost of the solution involving matrix inversion is much reduced with the $N \times N$ matrix $\mathbf{C}_{\bar{\mathbf{y}} \bar{\mathbf{y}}}^{-1}$ of Eqs. (8) and (9) in comparison with the $MN \times MN$ matrix $\mathbf{C}_{\mathbf{z} \mathbf{z}}^{-1}$ in Eq. (6).

From Eqs. (B.1)–(B.4), it is interesting to note that the node outputs $\bar{y}_n$ and the calculations of N diagonal elements of $\mathbf{C}_{\bar{\mathbf{y}} \bar{\mathbf{y}}}$ involve the sensor number M . However, the correlation matrix $\mathbf{C}_{\boldsymbol{\theta} \bar{\mathbf{y}}}$ and

the $N(N-1)$ non-diagonal elements of $\mathbf{C}_{\bar{\mathbf{y}}\bar{\mathbf{y}}}$ are unrelated to the sensor number M . In the following, it will be seen that the designed vector estimator $\hat{\boldsymbol{\theta}}_{\text{NE}}^M$ in Eq. (10) can be further simplified.

Corollary 2. For a sufficiently large number M of identical transfer functions $g_n(\cdot)$ in each network node, the designed estimator $\hat{\boldsymbol{\theta}}_i^M$ of Eq. (8) has an explicit form

$$\hat{\boldsymbol{\theta}}_i^* = \mathbf{E}_{\theta_i}(\theta_i) + \mathbf{C}_{\theta_i \mathbf{g}^*} \mathbf{C}_{\mathbf{g}^* \mathbf{g}^*}^{-1} \{\mathbf{g}^*(\mathbf{x}) - \mathbf{E}_{\mathbf{x}}[\mathbf{g}^*(\mathbf{x})]\}, \quad (12)$$

where the network output vector $\mathbf{g}^*(\mathbf{x}) = [g_1^*(x_1), g_2^*(x_2), \dots, g_N^*(x_N)]^\top$ and $g_n^*(x_n) = \mathbf{E}_{\eta}[g_n(x_n + \eta)]$ for $n = 1, 2, \dots, N$. The Bayesian MSE of the estimator $\hat{\boldsymbol{\theta}}_i^*$ of Eq. (12) can be also computed as

$$B_{\hat{\boldsymbol{\theta}}_i^*} = \mathbf{C}_{\theta_i} - \mathbf{C}_{\theta_i \mathbf{g}^*} \mathbf{C}_{\mathbf{g}^* \mathbf{g}^*}^{-1} \mathbf{C}_{\mathbf{g}^* \theta_i}, \quad (13)$$

where the $1 \times N$ correlation vector $\mathbf{C}_{\theta_i \mathbf{g}^*} = \mathbf{E}_{\mathbf{x}}[(\theta_i - \mathbf{E}_{\theta_i}(\theta_i))(\mathbf{g}^*(\mathbf{x}) - \mathbf{E}_{\mathbf{x}}[\mathbf{g}^*(\mathbf{x})])^\top]$ and the $N \times N$ covariance matrix $\mathbf{C}_{\mathbf{g}^* \mathbf{g}^*} = \mathbf{E}_{\mathbf{x}}\{(\mathbf{g}^*(\mathbf{x}) - \mathbf{E}_{\mathbf{x}}[\mathbf{g}^*(\mathbf{x})])(\mathbf{g}^*(\mathbf{x}) - \mathbf{E}_{\mathbf{x}}[\mathbf{g}^*(\mathbf{x})])^\top\}$. The vector estimator is given by

$$\begin{aligned} \hat{\boldsymbol{\theta}}_{\text{NE}}^* &= [\hat{\theta}_1^*, \hat{\theta}_2^*, \dots, \hat{\theta}_p^*]^\top \\ &= \mathbf{E}_{\theta}(\boldsymbol{\theta}) + \mathbf{C}_{\theta \mathbf{g}^*} \mathbf{C}_{\mathbf{g}^* \mathbf{g}^*}^{-1} \{\mathbf{g}^*(\mathbf{x}) - \mathbf{E}_{\mathbf{x}}[\mathbf{g}^*(\mathbf{x})]\}, \end{aligned} \quad (14)$$

where the $p \times N$ correlation matrix $\mathbf{C}_{\theta \mathbf{g}^*} = \mathbf{C}_{\mathbf{g}^* \theta}^\top = [\mathbf{C}_{\theta_1 \mathbf{g}^*}^\top, \mathbf{C}_{\theta_2 \mathbf{g}^*}^\top, \dots, \mathbf{C}_{\theta_p \mathbf{g}^*}^\top]^\top$. The Bayesian matrix of the vector estimator $\hat{\boldsymbol{\theta}}_{\text{NE}}^*$ is given by

$$\mathbf{B}_{\text{NE}}^* = \mathbf{E}_{\mathbf{x}}[(\boldsymbol{\theta} - \hat{\boldsymbol{\theta}}_{\text{NE}}^*)(\boldsymbol{\theta} - \hat{\boldsymbol{\theta}}_{\text{NE}}^*)^\top] = \mathbf{C}_{\theta\theta} - \mathbf{C}_{\theta \mathbf{g}^*} \mathbf{C}_{\mathbf{g}^* \mathbf{g}^*}^{-1} \mathbf{C}_{\mathbf{g}^* \theta} \quad (15)$$

whose diagonal element $[\mathbf{B}_{\text{NE}}^*]_{ii} = B_{\hat{\theta}_i^*}$ for $i = 1, 2, \dots, p$.

Proof of Corollary 2 is presented in Appendix C. It will be seen in Fig. 7 that, for given background noise and added noise, the designed vector estimator $\hat{\boldsymbol{\theta}}_{\text{NE}}^*$ in Eq. (14) can achieve the minimum Bayesian MSE that sets the lower MSE bound for $\hat{\boldsymbol{\theta}}_{\text{NE}}^M$ in Eq. (8).

Furthermore, in order to improve the performance of $\hat{\boldsymbol{\theta}}_{\text{NE}}^M$ including $\hat{\boldsymbol{\theta}}_{\text{NE}}^*$ for $M \rightarrow \infty$ by optimally tuning and exploiting the added noise, we face the minimization problem of the trace of the Bayesian matrix, i.e.

$$f_{\eta}^o(\eta) = \arg \min_{f_{\eta}} \text{tr}(\mathbf{B}_{\text{NE}}^M), \quad (16)$$

to find the optimal added noise PDF $f_{\eta}^o(\eta)$. This functional optimization problem of Eq. (16) is usually analytically intractable. For a numerical resolution, we use a kernel method [32,34] to find an approximate optimal solution as

$$\tilde{f}_{\eta}^o(\eta) = \sum_{k=1}^K \lambda_k \phi(\eta, \mu_k, \sigma_k), \quad (17)$$

where the normalization coefficients $\lambda_k \geq 0$ satisfy the constraint $\sum_{k=1}^K \lambda_k = 1$, and the Gaussian kernel function $\phi(\mu) = \exp[-(\mu - \mu_k)^2 / 2\sigma_k^2] / \sqrt{2\pi\sigma_k^2}$ with parameters μ_k and $\sigma_k \geq 0$. Moreover, it has been theoretically proved that, as the number of kernel functions K increases in Eq. (17), the approximate optimal noise PDF $\tilde{f}_{\eta}^{\text{opt}}$ can converge to the optimal noise PDF f_{η}^{opt} if it exists [32,43,44]. Thus, for a given kernel function number K , the minimization problem of Eq. (16) reduces to a finite-dimensional nonlinear constrained optimization

$$\begin{aligned} \min_{\lambda_k, \mu_k, \sigma_k} \quad & \text{tr}(\mathbf{B}_{\text{NE}}^M) \\ \text{s.t.} \quad & \lambda_k \geq 0, \sum_{k=1}^K \lambda_k = 1, \sigma_k \geq 0, \end{aligned} \quad (18)$$

with respect to parameters λ_k , μ_k and σ_k for $k = 1, 2, \dots, K$. The sequential quadratic programming method [44] has been implemented to solve this nonlinear constrained optimization of Eq. (18). At each major iteration, an approximation is made of the Hessian of the Lagrangian function using a quasi-Newton (BFGS) updating method [44]. This is then used to generate a quadratic programming subproblem whose solution is used to form a search direction for a line search procedure [44].

In the following parts, we will present some illustrative examples to show improvement of the designed estimator $\hat{\boldsymbol{\theta}}_{\text{NE}}^M$ by the purposeful addition of artificial noise. For reference, we consider the case where the parameter vector $\boldsymbol{\theta}$ is directly estimated based on the measurement vector $\mathbf{x}$, by the LMMSE unbiased estimator given by [39,40]

$$\hat{\boldsymbol{\theta}}_{\text{LMMSE}} = \mathbf{E}_{\theta}(\boldsymbol{\theta}) + \mathbf{C}_{\theta \mathbf{x}} \mathbf{C}_{\mathbf{x} \mathbf{x}}^{-1} [\mathbf{x} - \mathbf{E}_{\mathbf{x}}(\mathbf{x})], \quad (19)$$

with its Bayesian MSE matrix [39,40]

$$\begin{aligned} \mathbf{B}_{\text{LMMSE}} &= \mathbf{E}_{\mathbf{x}}[(\boldsymbol{\theta} - \hat{\boldsymbol{\theta}}_{\text{LMMSE}})(\boldsymbol{\theta} - \hat{\boldsymbol{\theta}}_{\text{LMMSE}})^\top] \\ &= \mathbf{C}_{\theta\theta} - \mathbf{C}_{\theta \mathbf{x}} \mathbf{C}_{\mathbf{x} \mathbf{x}}^{-1} \mathbf{C}_{\mathbf{x} \theta}^\top, \end{aligned} \quad (20)$$

where the vector mean $\mathbf{E}_{\mathbf{x}}(\mathbf{x}) = \mathbf{H} \mathbf{E}_{\theta}(\boldsymbol{\theta})$, the cross-covariance matrix $\mathbf{C}_{\theta \mathbf{x}} = \mathbf{C}_{\theta\theta} \mathbf{H}^\top$, the covariance matrix $\mathbf{C}_{\mathbf{x} \mathbf{x}} = \mathbf{H} \mathbf{C}_{\theta\theta} \mathbf{H}^\top + \mathbf{C}_{\xi\xi}$ and the covariance matrix $\mathbf{C}_{\xi\xi}$ of the background noise vector ξ are given a priori. Note that the LMMSE estimator $\hat{\boldsymbol{\theta}}_{\text{LMMSE}}$ in Eq. (19) is identical in form to the MMSE estimator for jointly Gaussian distribution of parameters and background noise [39,40].

3. Results

We will first discuss some illustrative examples to show the improvement of $\hat{\boldsymbol{\theta}}_{\text{NE}}^*$ of Eq. (14) by the purposeful addition of optimal noise η_{nm} , wherein the measurement vector $\mathbf{x}$ has a joint Gaussian distribution.

Example 1. Consider a random vector parameter $\boldsymbol{\theta} = [\theta_1, \theta_2]^\top$ in the linear vector model. Parameters θ_1 and θ_2 are independent and follow the Gaussian PDF $f_{\theta_i}(\theta_i) = \exp[-(\theta_i - u_{\theta_i})^2 / 2\sigma_{\theta_i}^2] / \sqrt{2\pi\sigma_{\theta_i}^2}$ with variance $\sigma_{\theta_i}^2$ for $i = 1, 2$, respectively. Assume that the vector mean $\mathbf{E}_{\theta}(\boldsymbol{\theta}) = [0, 0]^\top$ and the covariance matrix $\mathbf{C}_{\theta\theta} = [0.5, 0; 0, 0.1]$, where the semicolon denotes the end of a row of matrix. The zero-mean Gaussian background vector noise ξ is with the covariance matrix $\mathbf{C}_{\xi\xi} = 0.1 \mathbf{I}_N$ and the known $N \times 2$ observation matrix

$$\mathbf{H} = \begin{bmatrix} 1 & 0 \\ \cos(\frac{2\pi}{N}) & \sin(\frac{2\pi}{N}) \\ \vdots & \vdots \\ \cos(\frac{2\pi(N-1)}{N}) & \sin(\frac{2\pi(N-1)}{N}) \end{bmatrix}.$$

Consider all network nodes have the identical sensor transfer function

$$g(u) = \begin{cases} 1, & u > \gamma, \\ 0, & u \leq \gamma, \end{cases} \quad (21)$$

which is also known as a binary quantizer with the threshold γ [7,18,19,29,39].

Without the added noise η_{nm} , the node output of M identical binary quantizers in each node equals to an arbitrary quantizer output $g(x_n)$. Thus, the MSE matrix $\mathbf{B}_{\text{NE}}^*$ in Eq. (15) of the designed estimator $\hat{\boldsymbol{\theta}}_{\text{NE}}^*$ of Eq. (14) can only be optimized by the quantizer threshold γ . It is shown in Fig. 3 that, at the optimal threshold $\gamma = 0$ and without the added noise ($\sigma_{\eta} = 0$), the Bayesian

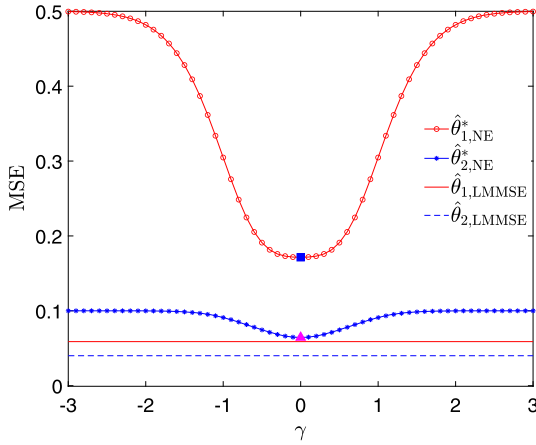


Fig. 3. Bayesian MSEs $[\mathbf{B}_{NE}^*]_{11}$ and $[\mathbf{B}_{NE}^*]_{22}$ of the designed estimator $\hat{\theta}_{NE}^* = [\hat{\theta}_{1,NE}^*, \hat{\theta}_{2,NE}^*]^T$ in Eq. (14) versus the quantizer threshold γ without the added noise. The number of measurements x_n is $N = 3$. For comparison, the Bayesian MSEs $[\mathbf{B}_{LMMSE}]_{11} = 0.0588$ and $[\mathbf{B}_{LMMSE}]_{22} = 0.0400$ of the LMMSE estimator $\hat{\theta}_{LMMSE} = [\hat{\theta}_{1,LMMSE}, \hat{\theta}_{2,LMMSE}]^T$ in Eq. (19) are also plotted by lines. (For interpretation of the colors in the figure(s), the reader is referred to the web version of this article.)

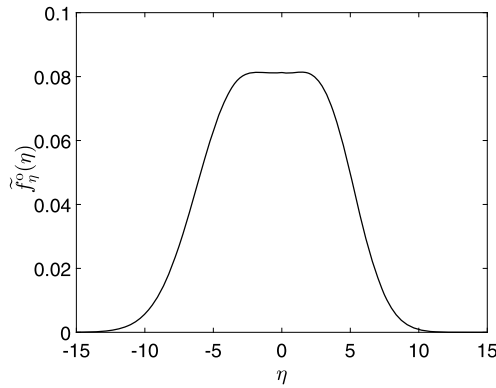


Fig. 4. Approximate optimal PDF $\tilde{f}_\eta^o(\eta)$ of the added noise for the noise-enhanced estimator $\hat{\theta}_{NE}^*$ with the kernel window number $K = 10$ in Eq. (17). The other parameters are the same as in Fig. 3.

MSEs of the estimator $\hat{\theta}_{NE}^* = [\hat{\theta}_{1,NE}^*, \hat{\theta}_{2,NE}^*]^T$ are minimized as $[\mathbf{B}_{NE}^*]_{11} = 0.1715$ (■) and $[\mathbf{B}_{NE}^*]_{22} = 0.0644$ (▲), respectively. Here, the number of measurements x_n is $N = 3$. For comparison, using Eq. (20), the Bayesian MSE matrix $\mathbf{B}_{LMMSE} = [0.0588, 0; 0, 0.0400]$ of the LMMSE estimator $\hat{\theta}_{LMMSE} = [\hat{\theta}_{1,LMMSE}, \hat{\theta}_{2,LMMSE}]^T$ can be calculated, and the corresponding Bayesian MSEs $[\mathbf{B}_{LMMSE}]_{11}$ and $[\mathbf{B}_{LMMSE}]_{22}$ are also plotted in Fig. 3 by lines. It is seen in Fig. 3 that, without the help of added noise, the performance of $\hat{\theta}_{NE}^*$ is significantly worse than that of $\hat{\theta}_{LMMSE}$, even when the threshold γ is optimized.

However, with the help of the optimal added noise, the much lower Bayesian MSE of the designed estimator can be expected. For a given window number $K = 10$, we use the sequential quadratic programming algorithm [44] to numerically solve the nonlinear constrained optimization problem in Eq. (18), and obtain the optimized parameters λ_k , μ_k and σ_k for $k = 1, 2, \dots, K$. Substituting these parameters into Eq. (17), the approximate optimal noise PDF $\tilde{f}_\eta^o$ can be established as illustrated in Fig. 4. Using this approximate optimal PDF $\tilde{f}_\eta^o$, we can calculate the corresponding Bayesian MSE matrix as $\mathbf{B}_{NE}^* = [0.0589, 0; 0, 0.0400]$, which is almost equal to the minimum MSE matrix $\mathbf{B}_{LMMSE} = [0.0588, 0; 0, 0.0400]$ obtained by the estimator $\hat{\theta}_{LMMSE}$ and far better than the Bayesian MSE matrix $\mathbf{B}_{NE}^* = [0.1715, 0; 0, 0.0644]$ obtained without the added noise. The distribution of noise com-

ponents η_{nm} according to the approximate optimal PDF $\tilde{f}_\eta^o(\eta)$ of Fig. 4 is nontrivial, and its beneficial role in the sensor network for vector parameter estimation is clearly manifested through greatly reducing the Bayesian MSEs of the designed estimator.

An interesting result is that the optimization of the Bayesian MSE matrix $\mathbf{B}_{NE}^*$ by the optimal noise is independent of the threshold value of γ . Whatever the quantizer threshold $\gamma = 1$ or 0 , the same minimum Bayesian MSE matrix $\mathbf{B}_{NE}^* = [0.0589, 0; 0, 0.0400]$ can be achieved by the corresponding optimal PDF solutions of $\tilde{f}_\eta^o(\eta)$. The reason is that the optimal PDF $\tilde{f}_\eta^o(\eta)$ contains the adjustable mean parameters μ_k that can eliminate the influence of the quantizer threshold on the parameter vector estimation, so that the performance of the designed estimator $\hat{\theta}_{NE}^*$ can be maintained at a stable and high level.

Another interesting problem is whether or not the designed estimator $\hat{\theta}_{NE}^*$ of Eq. (14) based on the binary data approaches the LMMSE estimator $\hat{\theta}_{LMMSE}$ of Eq. (19) that uses the complete analog data $\mathbf{x}$, no matter how the background noise intensity σ_ξ varies. In Fig. 5, we illustratively show the Bayesian MSEs of the estimator $\hat{\theta}_{NE}^*$ versus the background noise intensity $\sigma_\xi \in [0.1, 4]$ (the input SNR $\sigma_\theta^2/\sigma_\xi^2$). It is seen in Fig. 5 (a) and (b) that, for a wide range of input SNRs, the estimator $\hat{\theta}_{NE}^*$ in Eq. (14) improved by the optimal added noise has a better estimation performance than the case without added noise. Especially, it is well known that, for the jointly Gaussian vector parameter θ and background noise ξ , the LMMSE estimator $\hat{\theta}_{LMMSE}$ is just the MMSE estimator [39]. As illustrated in Fig. 5, for both parameters θ_1 and θ_2 , the Bayesian MSEs of $\hat{\theta}_{NE}^*$ are very close to the minimum Bayesian MSEs achieved by $\hat{\theta}_{LMMSE}$ over the considered range of input SNRs.

Example 2. As the observation number N of x_n increases, the Bayesian MSEs of the noise-enhanced estimator $\hat{\theta}_{NE}^*$ of Eq. (14) is illustrated in Fig. 6. With the help of the optimal added noise, the Bayesian MSEs of $\hat{\theta}_{NE}^*$ all nearly approach the minimum Bayesian MSEs achieved by $\hat{\theta}_{LMMSE}$ for different observation numbers.

It is noted that, based on the M-level uniformly quantized data, the trace of the Bayesian MSE matrix of the LMMSE estimator is much higher than that of the BCRLB matrix at the low power allocation [3]. It is known that the BCRLB is just the Bayesian MSE achieved by the LMMSE estimator $\hat{\theta}_{LMMSE}$ in the circumstance of the joint Gaussian distribution of parameters and background noise. In Figs. 5 and 6, based on binary quantized data, the Bayesian MSE of the designed estimator $\hat{\theta}_{NE}^*$ in Eq. (14) improved by the optimal added noise is very close to the BCRLB indicated by the primary data $\mathbf{x}$. Therefore, compared with the LMMSE estimator [3,45] without exploiting the added noise, the designed estimator $\hat{\theta}_{NE}^*$ in Eq. (14) has improved performance when the artificially added noise is optimized.

Example 3. In Example 1, the sensor number M in each sensor network node is assumed to be sufficiently large. However, large sensor number M may not be practical to implement. An important question is how sufficiently large the sensor number M in each node needs to be as the performance of the designed estimator $\hat{\theta}_{NE}^M$ of Eq. (8) approaches that of the LMMSE estimator $\hat{\theta}_{LMMSE}$ of Eq. (19).

In Fig. 7, it is seen that the Bayesian MSEs of the designed estimator $\hat{\theta}_{NE}^M$ in Eq. (8) monotonically decrease as the sensor number M increases. When the sensor number $M \geq 3000$, it is shown in Fig. 7 that the designed estimator $\hat{\theta}_{NE}^M$ in Eq. (8), assisted by the optimal added noise, has the corresponding minimum Bayesian MSEs given by diagonal elements $[\mathbf{B}_{NE}^M]_{11} = 0.0598$ and $[\mathbf{B}_{NE}^M]_{22} = 0.0403$, which is very close to the minimum Bayesian MSEs $[\mathbf{B}_{LMMSE}]_{11} = 0.0588$ and $[\mathbf{B}_{LMMSE}]_{22} = 0.0400$ of the LMMSE estimator $\hat{\theta}_{LMMSE}$. For different background noise intensities $\sigma_\xi \in$

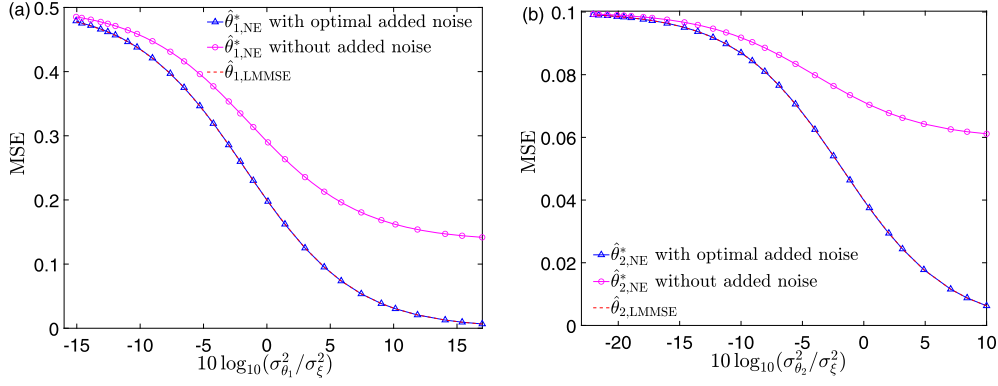


Fig. 5. Bayesian MSEs of the noise-enhanced estimator $\hat{\theta}_{NE}^*$ in Eq. (14) with the optimal added noise, $\hat{\theta}_{NE}^*$ without added noise and the LMMSE estimator $\hat{\theta}_{LMMSE}$ in Eq. (19) as a function of the input SNR $\sigma_{\theta}^2/\sigma_{\xi}^2$ for estimating (a) the parameter θ_1 and (b) the parameter θ_2 . Here, the parameters θ_1 and θ_2 are Gaussian distributed with the variances $\sigma_{\theta_1}^2 = 0.5$ and $\sigma_{\theta_2}^2 = 0.1$, respectively.

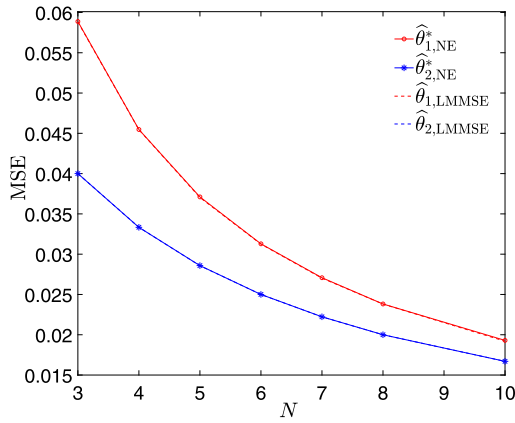


Fig. 6. Bayesian MSEs of the noise-enhanced estimator $\hat{\theta}_{NE}^*$ in Eq. (14) with the approximate optimal PDF $\tilde{f}_{\eta}^{opt}(\eta)$ versus the number N of the measurements x_n . The other parameters are the same as in Fig. 4.

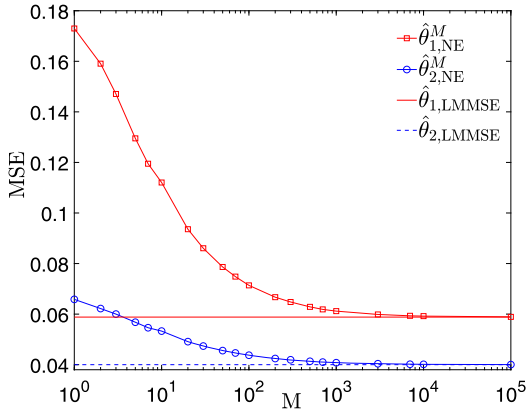


Fig. 7. Bayesian MSEs of the noise-enhanced estimator $\hat{\theta}_{NE}^M$ in Eq. (8) with the approximate optimal PDF $\tilde{f}_{\eta}^{opt}(\eta)$ versus the sensor number M in each network node. The other parameters are the same as in Fig. 4.

[0.1, 4], the Bayesian MSEs of the designed estimator $\hat{\theta}_{NE}^M$ are also observed to be comparable with that of $\hat{\theta}_{LMMSE}$ for a finite sensor number $M = 3000$ (not shown here).

It was proven [5] that, for a single-bit dithered quantizer and a Gaussian prior, any estimator using the binary quantizer output sequence is asymptotically at least $10\log(\pi/2) \approx 1.97$ dB worse than MMSE estimator using unquantized observations. In Fig. 7, it is seen that, for $M = 1$ (i.e. the single-bit dithered

quantizer), the Bayesian MSEs $[\mathbf{B}_{NE}^M]_{11}/[\mathbf{B}_{LMMSE}]_{11} > 1.57$ and $[\mathbf{B}_{NE}^M]_{22}/[\mathbf{B}_{LMMSE}]_{22} \approx 1.57$ also accord with the conclusion in [5]. However, upon the increasing the quantizer number M in each network node, it is shown in Fig. 7 that the designed estimator $\hat{\theta}_{NE}^M$ can break through this restriction of $10\log(\pi/2) \approx 1.97$ dB with the cost of establishing a large scale network. The reason is that, in the presence of the optimal added noise, a number of identical quantizers is demonstrated to be equivalent to an M-level quantizer [18,19].

Example 4. In Example 3, each network node has M identical sensors receiving the same weight coefficient. However, we may also investigate the impact of non-identical sensors which may arise in practice in a network. First, we consider $N = 3$ network nodes with quantizer thresholds $\gamma = 0, 0.2$, and 0.4 , respectively. The other parameters are the same as in Example 1. Using the optimal added noise indicated by the solved PDF $\tilde{f}_{\eta}^0$, the corresponding Bayesian MSE matrix $\mathbf{B}_{NE}^* = [0.0589, 0; 0, 0.0401]$ can be achieved. Compared with the Bayesian MSEs obtained in Example 1 with identical quantizers over all nodes, the designed estimator $\hat{\theta}_{NE}^*$ with non-identical quantizers performs almost the same. Secondly, we settle $M = 3000$ binary quantizers in each network node, but 50 percent of the quantizers is with the threshold $\gamma = 0$ and the other half is with $\gamma = 0.5$. With the injection of optimal added noise, the Bayesian MSE matrix is obtained as $\mathbf{B}_{NE}^M = [0.0605, 0; 0, 0.0410]$. Compared with the Bayesian MSEs $[\mathbf{B}_{NE}^M]_{11} = 0.0598$ and $[\mathbf{B}_{NE}^M]_{22} = 0.0403$ obtained by $\hat{\theta}_{NE}^M$ with $M = 3000$ identical quantizers in Example 3, the 2% performance degradation of $\hat{\theta}_{NE}^M$ with non-identical sensors is at tolerable cost.

In Examples 1–4, using the property of multivariate Gaussian PDF, the minimization problem of the trace of the Bayesian MSE matrix in Eq (18) by the optimal noise becomes possible. However, some parameters to be estimated may have a non-Gaussian distribution in practice. As a consequence, the characterization of the MSE via the second-order moments of Eqs. (B.2)–(B.4) cannot be performed analytically, and a numerical approach has to be implemented. In the next example, we consider another practical vector parameter estimation via numerically computing the first two moments of the sensor outputs, where the added Gaussian noise level is numerically optimized to improve the performance of the designed estimator $\hat{\theta}_{NE}^M$ in Eq. (8).

Example 5. Let mutually independent random parameters θ_1 and θ_2 uniformly distribute in the intervals $[0, 1]$ and $[0, 2]$, respectively. The sensor transfer function is a hard-limiting function

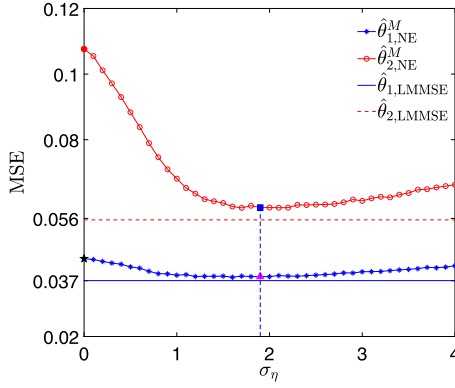


Fig. 8. Plots of Bayesian MSEs of the noise-enhanced estimator $\hat{\theta}_{NE}^M$ as a function of the added Gaussian noise level σ_η . Here, the random parameters θ_1 and θ_2 uniformly distribute in the intervals $[0, 1]$ and $[0, 2]$, respectively. The sensor number $M = 10^3$ in each node, and the number of simulation trials is 10^5 .

$$g(u) = \begin{cases} 0, & u < 0, \\ u, & 0 \leq u < 1, \\ 1, & u \geq 1, \end{cases} \quad (22)$$

which is also the maximum *a posteriori* estimator for estimating a scalar uniform parameter buried in Gaussian noise [32,39,40]. The size of the measurement vector $\mathbf{x}$ is $N = 3$, and the zero-mean background noise ξ has the covariance matrix $\mathbf{C}_\xi = 0.1\mathbf{I}_3$. For a given added noise level σ_η , we add the zero-mean Gaussian noise components η_{nm} to network nodes. Using the Algorithm 1, the first-order and second-order moments of Eqs. (B.2)–(B.4) are numerically calculated, and the designed estimator $\hat{\theta}_{NE}^M$ in Eq. (10) and its Bayesian MSE matrix $\mathbf{B}_{NE}^M$ of Eq. (11) are also obtained.

Using the Algorithm 1 and without the added noise ($\sigma_\eta = 0$), the Bayesian MSE matrix of the designed estimator $\hat{\theta}_{NE}^M$ is numerically computed as $\mathbf{B}_{NE}^M = [0.0434, 0.0141; 0.0141, 0.1076]$. Here, the sensor number $M = 10^3$ in each network node. Therefore, as illustrated in Fig. 8, the Bayesian MSEs of $[\mathbf{B}_{NE}^M]_{11}$ and $[\mathbf{B}_{NE}^M]_{22}$ start from 0.0434 and 0.1076 at $\sigma_\eta = 0$, respectively. It is shown in Fig. 8 that, upon increasing the added noise level σ_η , the Bayesian MSE matrix can be minimized as $\mathbf{B}_{NE}^M = [0.0382, 0.0008; 0.0008, 0.0593]$ at the optimal added noise level $\sigma_\eta^{\text{opt}} = 1.9$. Compared to the case of no added noise ($\sigma_\eta = 0$), the MSE of $\hat{\theta}_{1,NE}^M$ is reduced from 0.0434 (★) to 0.0382 (▲), and the MSE of $\hat{\theta}_{2,NE}^M$ can be reduced from 0.1076 (●) to 0.0593 (■), as shown in Fig. 8. For comparison, the Bayesian MSEs of $\hat{\theta}_{1,LMMSE}$ and $\hat{\theta}_{2,LMMSE}$ are also illustrated in Fig. 8 by lines, and the Bayesian MSE matrix is $\mathbf{B}_{LMMSE}^M = [0.0370, 0; 0, 0.0556]$ achieved by the LMMSE estimator $\hat{\theta}_{LMMSE}^M$ based on the complete analog data $\mathbf{x}$.

Furthermore, for different input SNRs of $\sigma_\theta^2/\sigma_\xi^2$, the corresponding Bayesian MSEs upon increasing the noise intensity σ_ξ of the background noise ξ are also illustrated in Fig. 9. It is seen in Fig. 9 (a) and (b) that, for the considered range of the background noise intensity σ_ξ , the designed estimator $\hat{\theta}_{NE}^M$ aided by the optimal added Gaussian noise does outperform the one without the added noise, and is also very close to the LMMSE estimator $\hat{\theta}_{LMMSE}^M$ based on the complete analog data $\mathbf{x}$ at low input SNRs of $\sigma_\theta^2/\sigma_\xi^2$. Note that there still is a visible difference between the MSEs of the noise-enhanced estimator $\hat{\theta}_{NE}^M$ and the LMMSE estimator $\hat{\theta}_{LMMSE}^M$ at high input SNRs of $\sigma_\theta^2/\sigma_\xi^2$, because the added noise type is restricted to a given Gaussian distribution and the only parameter to be optimized is the noise level σ_η under the considered conditions in Example 5.

Input: Data number N , sensor number M in every node, sensor function g , parameter vectors θ^t with known mean vector $E_\theta(\theta)$, background noise vectors ξ^t , measurement vectors $\mathbf{x}^t = \mathbf{H}\theta^t + \xi^t$ and the added noise vectors η_n^t with a given noise level σ_η of added Gaussian noise. t is the times of simulation and T is the total simulation number.

Output: $\hat{\theta}_{NE}^M, \mathbf{B}_{NE}^M$.

```

1 for  $n=1 \dots N$  do
2   for  $t=1 \dots T$  do
3      $\bar{\mathbf{y}}_n^t = \frac{1}{M} \sum_{m=1}^M g(x_n^t \mathbf{1}_M + \eta_n^t)$ ;
4   end
5    $E(\bar{\mathbf{y}}_n) \approx \frac{1}{T} \sum_{t=1}^T \bar{\mathbf{y}}_n^t$ ;
6 end
7 for  $t=1 \dots T$  do
8   for  $n=1 \dots N$  do
9     for  $i=1 \dots p$  do
10       $[\mathbf{C}_{\theta\bar{\mathbf{y}}}]_{in}^t = [\theta_i^t - E_{\theta_i}(\theta_i)][\bar{\mathbf{y}}_n^t - E(\bar{\mathbf{y}}_n)]$ ;
11    end
12    for  $\ell=n \dots N$  do
13       $[\mathbf{C}_{\bar{\mathbf{y}}\bar{\mathbf{y}}}]_{n\ell}^t = [\bar{\mathbf{y}}_n^t - E(\bar{\mathbf{y}}_n)][\bar{\mathbf{y}}_\ell^t - E(\bar{\mathbf{y}}_\ell)]$ ;
14    end
15  end
16   $[\mathbf{C}_{\theta\bar{\mathbf{y}}}]_{in} = \frac{1}{T} \sum_{t=1}^T [\mathbf{C}_{\theta\bar{\mathbf{y}}}]_{in}^t$ ;
17   $[\mathbf{C}_{\bar{\mathbf{y}}\bar{\mathbf{y}}}]_{n\ell} = \frac{1}{T} \sum_{t=1}^T [\mathbf{C}_{\bar{\mathbf{y}}\bar{\mathbf{y}}}]_{n\ell}^t$ ;
18   $\hat{\theta}_{NE}^{M,t} \approx E_\theta(\theta) + \mathbf{C}_{\theta\bar{\mathbf{y}}} \mathbf{C}_{\bar{\mathbf{y}}\bar{\mathbf{y}}}^{-1} \begin{pmatrix} \bar{\mathbf{y}}_1^t - E(\bar{\mathbf{y}}_1) \\ \bar{\mathbf{y}}_2^t - E(\bar{\mathbf{y}}_2) \\ \vdots \\ \bar{\mathbf{y}}_N^t - E(\bar{\mathbf{y}}_N) \end{pmatrix}$ ;
19 end
20  $\mathbf{B}_{NE}^M \approx \frac{1}{T} \sum_{t=1}^T (\theta^t - \hat{\theta}_{NE}^{M,t})(\theta^t - \hat{\theta}_{NE}^{M,t})^\top$ .
```

Algorithm 1: Numerical calculations of $\hat{\theta}_{NE}^M$ of Eq. (10) and $\mathbf{B}_{NE}^M$ of Eq. (11).

4. Conclusion

In this paper, we design a noise-enhanced Bayesian vector estimator based on low-resolution observations transmitted by sensor networks, where each network node is composed of a large number of low-complexity sensors. When the sensors in the network nodes are identical, we present a simplified calculation method for the design of the noise-enhanced estimator and its the Bayesian MSE matrix by reducing the dimensions of the weight vector and network output vector. For illustrative examples with a joint Gaussian distribution of the vector parameter and background noise, we obtain a characterization of the optimal added noise that provides a Bayesian MSE comparable with that of the MMSE estimator built on the primary data. For a joint non-Gaussian case, we intentionally add an optimal amount of Gaussian noise into the designed estimator and also obtain much reduced Bayesian MSE by numerical experiments. The derivations and examples illustrate the feasibility of exploiting added noise in low-complexity sensor networks to enhance the efficiency for information processing.

Some open questions still need to be addressed in future studies. For instance, when the joint PDF of the vector parameter and background noise is non-Gaussian, the theoretical solution of the optimal added noise remains unsolved. We set the type of added noise and only optimize the added noise level to minimize the Bayesian MSE of the designed estimator. Thus, the difference between the MSEs of the noise-enhanced estimator $\hat{\theta}_{NE}^M$ and the LMMSE estimator $\hat{\theta}_{LMMSE}^M$ is evident at high input SNRs, as shown in Fig. 9. How to further reduce the MSE of the noise-enhanced estimator $\hat{\theta}_{NE}^M$ in the considered case of the joint non-Gaussian PDF of vector parameter and background noise is an open question for future study. In addition, the evaluations of the designed estimator

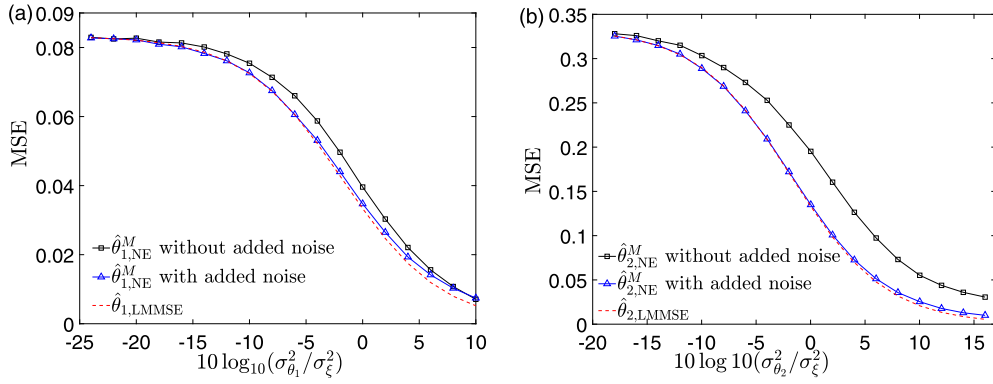


Fig. 9. Bayesian MSEs of the designed noise-enhanced estimator $\hat{\theta}_{NE}^*$ in Eq. (11) with the optimal added Gaussian noise, $\hat{\theta}_{NE}^M$ without added noise and the LMMSE estimator $\hat{\theta}_{LMMSE}$ in Eq. (19) as a function of the input SNR $\sigma_\theta^2/\sigma_\xi^2$ for (a) the parameter θ_1 and (b) the parameter θ_2 . The other parameters are the same as in Fig. 8.

and its Bayesian MSE still cannot avoid the matrix inversion, although the matrix dimension is greatly reduced for each network node consisting of identical sensors. Alternatively, in other signal estimation problems the observed data may be on-going in that as time progresses, more data become available and need to be processed sequentially in time [39]. Then, the sequential implementation of the noise-enhanced estimator with no matrix inversion will also be of great interest.

CRediT authorship contribution statement

Fabing Duan: Conceptualization, Methodology, Creation of models. **Jia Liu:** Data curation, Programming, Writing - Original draft preparation. **François Chapeau-Blondeau and Derek Abbott:** Investigation, Validation, Writing - Reviewing and Editing.

Declaration of competing interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Acknowledgments

This work is supported in part by the Taishan Scholar Project of Shandong Province of China (No. ts20190930) and Shandong Key Laboratory of Industrial Control Technology.

Appendix A. Proof of Theorem 1

Without loss of generality, we set M identical transfer functions $g_{1m}(\cdot) = g_1(\cdot)$ of sensors in the first network node. Letting $\omega = [\omega_1^\top, \tilde{\omega}^\top]^\top$ ($\tilde{\omega} = [\omega_2^\top, \omega_3^\top, \dots, \omega_N^\top]^\top$) and $\mathbf{y} = [\mathbf{y}_1^\top, \tilde{\mathbf{y}}^\top]^\top$ ($\tilde{\mathbf{y}} = [\mathbf{y}_2^\top, \mathbf{y}_3^\top, \dots, \mathbf{y}_N^\top]^\top$), the covariance matrix $\mathbf{C}_{zz}$ can be partitioned as

$$\mathbf{C}_{zz} = \mathbf{E}_{\mathbf{x},\eta} \{ [\mathbf{z} - \mathbf{E}_{\mathbf{x},\eta}(\mathbf{z})][\mathbf{z} - \mathbf{E}_{\mathbf{x},\eta}(\mathbf{z})]^\top \} = \begin{bmatrix} \mathbf{C}_{\mathbf{y}_1\mathbf{y}_1} & \mathbf{C}_{\mathbf{y}_1\tilde{\mathbf{y}}} \\ \mathbf{C}_{\tilde{\mathbf{y}}\mathbf{y}_1} & \mathbf{C}_{\tilde{\mathbf{y}}\tilde{\mathbf{y}}} \end{bmatrix}, \quad (\text{A.1})$$

where the partitioned matrices $\mathbf{C}_{\mathbf{y}_1\mathbf{y}_1} = \mathbf{E}_{\mathbf{x}_1,\eta} \{ [\mathbf{y}_1 - \mathbf{E}_{\mathbf{x}_1,\eta}(\mathbf{y}_1)][\mathbf{y}_1 - \mathbf{E}_{\mathbf{x}_1,\eta}(\mathbf{y}_1)]^\top \}$, $\mathbf{C}_{\mathbf{y}_1\tilde{\mathbf{y}}} = \mathbf{C}_{\tilde{\mathbf{y}}\mathbf{y}_1}^\top = \mathbf{E}_{\mathbf{x},\eta} \{ [\mathbf{y}_1 - \mathbf{E}_{\mathbf{x}_1,\eta}(\mathbf{y}_1)][\tilde{\mathbf{y}} - \mathbf{E}_{\mathbf{x},\eta}(\tilde{\mathbf{y}})]^\top \}$ and $\mathbf{C}_{\tilde{\mathbf{y}}\tilde{\mathbf{y}}} = \mathbf{E}_{\mathbf{x},\eta} \{ [\tilde{\mathbf{y}} - \mathbf{E}_{\mathbf{x},\eta}(\tilde{\mathbf{y}})][\tilde{\mathbf{y}} - \mathbf{E}_{\mathbf{x},\eta}(\tilde{\mathbf{y}})]^\top \}$. Similarly, the correlation vector $\mathbf{C}_{z\theta_i}$ can be partitioned as

$$\mathbf{C}_{z\theta_i} = [\mathbf{C}_{\mathbf{y}_1\theta_i}^\top, \mathbf{C}_{\tilde{\mathbf{y}}\theta_i}^\top]^\top \quad (\text{A.2})$$

with the partitioned correlation vector $\mathbf{C}_{\mathbf{y}_1\theta_i} = \mathbf{E}_{\mathbf{x}_1,\eta} \{ [\mathbf{y}_1 - \mathbf{E}_{\mathbf{x}_1,\eta}(\mathbf{y}_1)][\theta_i - \mathbf{E}_{\theta_i}(\theta_i)] \}$ and $\mathbf{C}_{\tilde{\mathbf{y}}\theta_i} = \mathbf{E}_{\mathbf{x},\eta} \{ [\tilde{\mathbf{y}} - \mathbf{E}_{\mathbf{x},\eta}(\tilde{\mathbf{y}})][\theta_i - \mathbf{E}_{\theta_i}(\theta_i)] \}$.

Since the optimum weight vector $\omega^{\text{opt}} = \mathbf{C}_{zz}^{-1} \mathbf{C}_{z\theta_i}$ can be rewritten as

$$\begin{bmatrix} \mathbf{C}_{\mathbf{y}_1\mathbf{y}_1} & \mathbf{C}_{\mathbf{y}_1\tilde{\mathbf{y}}} \\ \mathbf{C}_{\tilde{\mathbf{y}}\mathbf{y}_1} & \mathbf{C}_{\tilde{\mathbf{y}}\tilde{\mathbf{y}}} \end{bmatrix} \begin{bmatrix} \omega_1 \\ \tilde{\omega} \end{bmatrix} = \begin{bmatrix} \mathbf{C}_{\mathbf{y}_1\theta_i} \\ \mathbf{C}_{\tilde{\mathbf{y}}\theta_i} \end{bmatrix} \quad (\text{A.3})$$

thus we find

$$\mathbf{C}_{\mathbf{y}_1\mathbf{y}_1} \omega_1 + \mathbf{C}_{\mathbf{y}_1\tilde{\mathbf{y}}} \tilde{\omega} = \mathbf{C}_{\mathbf{y}_1\theta_i}. \quad (\text{A.4})$$

For M identical transfer functions $g_{1m}(\cdot) = g_1(\cdot)$ in the first network node and the identically distributed noise components η_{1m} , the covariance matrix $\mathbf{C}_{\mathbf{y}_1\mathbf{y}_1}$ has M equivalent diagonal elements

$$[\mathbf{C}_{\mathbf{y}_1\mathbf{y}_1}]_{mm} = \mathbf{E}_{\mathbf{x}_1,\eta} [g_1^2(x_1 + \eta)] - \mathbf{E}_{\mathbf{x}_1,\eta}^2 [g_1(x_1 + \eta)] \quad (\text{A.5})$$

and $M(M-1)$ equivalent non-diagonal elements

$$[\mathbf{C}_{\mathbf{y}_1\mathbf{y}_1}]_{mj} = \mathbf{E}_{\mathbf{x}_1} [\mathbf{E}_\eta^2 [g_1(x_1 + \eta)]] - \mathbf{E}_{\mathbf{x}_1,\eta}^2 [g_1(x_1 + \eta_n)] \quad (\text{A.6})$$

for $m, j = 1, 2, \dots, M$ ($m \neq j$). Its inverse matrix $\mathbf{C}_{\mathbf{y}_1\mathbf{y}_1}^{-1}$ also has M equivalent diagonal elements $[\mathbf{C}_{\mathbf{y}_1\mathbf{y}_1}^{-1}]_{mm}$ and $M(M-1)$ equivalent non-diagonal elements $[\mathbf{C}_{\mathbf{y}_1\mathbf{y}_1}^{-1}]_{mj}$. Therefore, the equality $\mathbf{1}_M^\top \mathbf{C}_{\mathbf{y}_1\mathbf{y}_1} \mathbf{C}_{\mathbf{y}_1\mathbf{y}_1}^{-1} \mathbf{1}_M = \mathbf{1}_M^\top \mathbf{1}_M = M$ yields

$$\begin{aligned} & [\mathbf{C}_{\mathbf{y}_1\mathbf{y}_1}^{-1}]_{mm} + (M-1)[\mathbf{C}_{\mathbf{y}_1\mathbf{y}_1}^{-1}]_{mj} \\ &= ([\mathbf{C}_{\mathbf{y}_1\mathbf{y}_1}]_{mm} + (M-1)[\mathbf{C}_{\mathbf{y}_1\mathbf{y}_1}]_{mj})^{-1}. \end{aligned} \quad (\text{A.7})$$

The correlation vector $\mathbf{C}_{\mathbf{y}_1\theta_i}$ also has M equivalent elements

$$\begin{aligned} [\mathbf{C}_{\mathbf{y}_1\theta_i}]_m &= \mathbf{E}_{\mathbf{x}_1,\eta} \{ \tilde{\mathbf{y}}_{1m} \tilde{\theta}_i \} \\ &= \mathbf{E}_{\mathbf{x}_1} \{ \theta_i \mathbf{E}_\eta [g_1(x_1 + \eta)] \} - \mathbf{E}_{\theta_i}(\theta_i) \mathbf{E}_{\mathbf{x}_1} \{ \mathbf{E}_\eta [g_1(x_1 + \eta)] \}. \end{aligned} \quad (\text{A.8})$$

Similarly, for M identical transfer functions $g_{1m}(\cdot) = g_1(\cdot)$, the $M \times (N-1)M$ partitioned matrix $\mathbf{C}_{\mathbf{y}_1\tilde{\mathbf{y}}}$ has M identical row vectors, which can be expressed as $\mathbf{d} = \mathbf{E}_{\mathbf{x},\eta} \{ [g_1(x_1 + \eta) - \mathbf{E}_{\mathbf{x}_1,\eta}(g_1(x_1 + \eta))] [\tilde{\mathbf{y}} - \mathbf{E}_{\mathbf{x},\eta}(\tilde{\mathbf{y}})]^\top \}$. Therefore, the vector $\mathbf{C}_{\mathbf{y}_1\tilde{\mathbf{y}}} \tilde{\omega}$ in Eq. (A.4) can be written as $\alpha \mathbf{1}_M$ with $\alpha = \mathbf{d}^\top \tilde{\omega}$.

Thus, from Eq. (A.4), we find $\mathbf{C}_{\mathbf{y}_1\mathbf{y}_1} \omega_1 = \mathbf{C}_{\mathbf{y}_1\theta_i} - \mathbf{C}_{\mathbf{y}_1\tilde{\mathbf{y}}} \tilde{\omega} = ([\mathbf{C}_{\mathbf{y}_1\theta_i}]_m - \alpha) \mathbf{1}_M$ and the optimum weight vector ω_1^{opt}

$$\omega_1^{\text{opt}} = \mathbf{C}_{\mathbf{y}_1\mathbf{y}_1}^{-1} ([\mathbf{C}_{\mathbf{y}_1\theta_i}]_m - \alpha) \mathbf{1}_M = w_1^{\text{opt}} \mathbf{1}_M, \quad (\text{A.9})$$

with

$$w_1^{\text{opt}} = \frac{[\mathbf{C}_{\mathbf{y}_1\theta_i}]_m - \alpha}{[\mathbf{C}_{\mathbf{y}_1\mathbf{y}_1}]_{mm} + (M-1)[\mathbf{C}_{\mathbf{y}_1\mathbf{y}_1}]_{mj}}. \quad (\text{A.10})$$

For any of network node with M identical sensors, Eq. (A.10) holds.

Appendix B. Proof of Corollary 1

For M identical transfer functions $g_{nm}(\cdot) = g_n(\cdot)$ of sensors in the n -th network node, each node has a scalar output $\bar{y}_n$ as shown in Fig. 2. Based on N outputs $\bar{y}_n$, the intermediate linear estimator can be expressed as $\hat{\theta}_i^M = w_0 + \mathbf{w}^\top \bar{\mathbf{y}}$, where w_0 is the biasing weight. In the same way, setting the derivatives $\partial E_{\mathbf{x},\eta}[(\theta_i - \hat{\theta}_i^M)^2]/\partial w_0 = 0$ and $\partial E_{\mathbf{x},\eta}[(\theta_i - \hat{\theta}_i^M)^2]/\partial \mathbf{w} = 0$, we find the optimum bias $w_0 = E_{\theta_i}(\theta_i) - \mathbf{w}^\top E_{\mathbf{x},\eta}(\bar{\mathbf{y}})$ and the optimal weight vector $\mathbf{w}^{\text{opt}} = \mathbf{C}_{\bar{\mathbf{y}}\bar{\mathbf{y}}}^{-1} \mathbf{C}_{\bar{\mathbf{y}}\theta}$. Substituting w_0 and $\mathbf{w}^{\text{opt}}$ into $\hat{\theta}_i^M = w_0 + \mathbf{w}^\top \bar{\mathbf{y}}$, we find the designed estimator $\hat{\theta}_i^M$ given by Eq. (8).

More specifically, the expectation of the node output $\bar{y}_n$ in Eq. (7) can be computed as

$$\begin{aligned} E_{x_n,\eta}(\bar{y}_n) &= \frac{1}{M} \sum_{m=1}^M E_{x_n,\eta}[g_n(x_n + \eta_{nm})] \\ &= E_{x_n}\{E_\eta[g_n(x_n + \eta)]\}, \end{aligned} \quad (\text{B.1})$$

because the added noise components η_{nm} have the common PDF f_η . Therefore, the correlation vector $\mathbf{C}_{\theta_i\bar{\mathbf{y}}} = \mathbf{C}_{\bar{\mathbf{y}}\theta_i}^\top = E_{\mathbf{x},\eta}\{[\theta_i - E_{\theta_i}(\theta_i)][\bar{\mathbf{y}} - E_{\mathbf{x},\eta}(\bar{\mathbf{y}})]^\top\}$ has elements

$$\begin{aligned} [\mathbf{C}_{\bar{\mathbf{y}}\theta_i}]_n &= E_{x_n,\eta}\{[\theta_i - E_{\theta_i}(\theta_i)][\bar{y}_n - E_{x_n,\eta}(\bar{y}_n)]\} \\ &= E_{x_n}\{E_\eta[g_n(x_n + \eta)]\} - E_{\theta_i}(\theta_i)E_{x_n}\{E_\eta[g_n(x_n + \eta)]\}, \end{aligned} \quad (\text{B.2})$$

which is independent of the sensor number M . The covariance matrix $\mathbf{C}_{\bar{\mathbf{y}}\bar{\mathbf{y}}}$ has N diagonal elements

$$\begin{aligned} [\mathbf{C}_{\bar{\mathbf{y}}\bar{\mathbf{y}}}]_{nn} &= E_{x_n,\eta}(\bar{y}_n^2) - E_{x_n,\eta}^2(\bar{y}_n) \\ &= E_{x_n,\eta}\left[\left(\frac{1}{M} \sum_{m=1}^M g_n(x_n + \eta_{nm})\right)^2\right] - E_{x_n,\eta}^2[g_n(x_n + \eta)] \\ &= \frac{1}{M^2} \left\{ \sum_{m=1}^M E_{x_n,\eta}[g_n^2(x_n + \eta_{nm})] \right. \\ &\quad \left. + E_{x_n}\left(\sum_{m=1}^M E_\eta[g_n(x_n + \eta_{nm})]\right) \right. \\ &\quad \left. \times \sum_{j=1}^M E_\eta[g_n(x_n + \eta_{nj})] \right\} (m \neq j) - E_{x_n,\eta}^2[g_n(x_n + \eta)] \\ &= \frac{1}{M} E_{x_n}\{E_\eta[g_n^2(x_n + \eta)]\} + \frac{M-1}{M} E_{x_n}\{E_\eta^2[g_n(x_n + \eta)]\} \\ &\quad - E_{x_n}^2\{E_\eta[g_n(x_n + \eta)]\}, \end{aligned} \quad (\text{B.3})$$

and $N(N-1)$ non-diagonal elements

$$\begin{aligned} [\mathbf{C}_{\bar{\mathbf{y}}\bar{\mathbf{y}}}]_{n\ell} &= E_{x_n,x_\ell,\eta}(\bar{y}_n \bar{y}_\ell) - E_{x_n,\eta}(\bar{y}_n)E_{x_\ell,\eta}(\bar{y}_\ell) (n \neq \ell) \\ &= \frac{1}{M^2} E_{x_n,x_\ell} \left(\sum_{m=1}^M E_\eta[g_n(x_n + \eta_{nm})] \sum_{j=1}^M E_\eta[g_\ell(x_\ell + \eta_{\ell j})] \right) \\ &\quad - E_{x_n,\eta}[g_n(x_n + \eta)]E_{x_\ell,\eta}[g_\ell(x_\ell + \eta)] \\ &= E_{x_n,x_\ell}\{E_\eta[g_n(x_n + \eta)]E_\eta[g_\ell(x_\ell + \eta)]\} \\ &\quad - E_{x_n}\{E_\eta[g_n(x_n + \eta)]\}E_{x_\ell}\{E_\eta[g_\ell(x_\ell + \eta)]\}. \end{aligned} \quad (\text{B.4})$$

It is noted that the $N(N-1)$ non-diagonal elements $[\mathbf{C}_{\bar{\mathbf{y}}\bar{\mathbf{y}}}]_{n\ell}$ are also independent of the sensor number M . With Eqs. (B.1)–(B.4), the p designed estimator $\hat{\theta}_i^M$ in Eq. (8) and its Bayesian MSEs in Eq. (9) can be obtained.

Appendix C. Proof of Corollary 2

As the sensor number $M \rightarrow \infty$ in N network nodes, the node output asymptotically converges to

$$\begin{aligned} \lim_{M \rightarrow \infty} \bar{y}_n &= \lim_{M \rightarrow \infty} \frac{1}{M} \sum_{m=1}^M g_n(x_n + \eta_{nm}) = E_\eta[g_n(x_n + \eta)] \\ &= g_n^*(x_n), \end{aligned} \quad (\text{C.1})$$

and the diagonal elements $[\mathbf{C}_{\bar{\mathbf{y}}\bar{\mathbf{y}}}]_{nn}$ in Eq. (B.3) of $\mathbf{C}_{\bar{\mathbf{y}}\bar{\mathbf{y}}}$ converge to

$$\begin{aligned} \lim_{M \rightarrow \infty} [\mathbf{C}_{\bar{\mathbf{y}}\bar{\mathbf{y}}}]_{nn} &= E_{x_n}\{E_\eta^2[g_n(x_n + \eta)]\} - E_{x_n}^2\{E_\eta[g_n(x_n + \eta)]\} \\ &= E_{x_n}[g_n^{*2}(x_n)] - E_{x_n}^2[g_n^*(x_n)] \\ &= [\mathbf{C}_{g^*g^*}]_{nn}, \end{aligned} \quad (\text{C.2})$$

because the term $E_{x_n}\{E_\eta[g_n^2(x_n + \eta)]\} < \infty$ in Eq. (B.3). It is also noted that, using Eq. (C.1), Eq. (B.2) and Eq. (B.4), we find that the elements $[\mathbf{C}_{\theta_i\bar{\mathbf{y}}}]_n = E_{x_n}\{E_\eta[g_n^*(x_n)]\} - E_{\theta_i}(\theta_i)E_{x_n}\{g_n^*(x_n)\} = [\mathbf{C}_{\theta g^*}]_n$ and the non-diagonal elements $[\mathbf{C}_{\bar{\mathbf{y}}\bar{\mathbf{y}}}]_{n\ell} = E_{x_n,x_\ell}\{g_n^*(x_n)g_\ell^*(x_\ell)\} - E_{x_n}\{g_n^*(x_n)\}E_{x_\ell}\{g_\ell^*(x_\ell)\} = [\mathbf{C}_{g^*g^*}]_{n\ell}$.

Substituting Eqs. (C.1) and (C.2) into Eq. (8), we have the explicit form of the estimator $\hat{\theta}_i^*$ in Eq. (12) and its Bayesian MSE $B_{\hat{\theta}_i^*}$ in Eq. (13). Combining p estimators of $\hat{\theta}_i^*$ in Eq. (12), we have the vector estimator $\hat{\boldsymbol{\theta}}_{\text{NE}}^*$ in Eq. (14) and its Bayesian MSE matrix $\mathbf{B}_{\text{NE}}^*$ in Eq. (15).

References

- [1] O. Dabeer, A. Karnik, Signal parameter estimation using 1-bit dithered quantization, *IEEE Trans. Inf. Theory* 52 (12) (2006) 5389–5405.
- [2] J. Xiao, R. Ribeiro, Z. Luo, G. Giannakis, Distributed compression-estimation using wireless sensor networks, *IEEE Signal Process. Mag.* 23 (4) (2006) 27–41.
- [3] M. Shieazi, A. Vosoughi, On Bayesian Fisher information maximization for distributed vector estimation, *IEEE Trans. Signal Inf. Process. Netw.* 5 (4) (2019) 628–645.
- [4] H. Chen, P.K. Varshney, Performance limit for distributed estimation systems with identical one-bit quantizers, *IEEE Trans. Signal Process.* 58 (1) (2010) 466–471.
- [5] G. Zeitler, G. Kramer, A. Singer, Bayesian parameter estimation using single-bit dithered quantization, *IEEE Trans. Signal Process.* 60 (6) (2012) 2713–2716.
- [6] F. Chapeau-Blondeau, S. Blanchard, D. Rousseau, Fisher information and noise-aided power estimation from one-bit quantizers, *Digit. Signal Process.* 18 (3) (2008) 434–443.
- [7] H.C. Papadopoulos, G.W. Wornell, A.V. Oppenheim, Sequential signal encoding from noisy measurements using quantizers with dynamic bias control, *IEEE Trans. Inf. Theory* 47 (3) (2001) 978–1002.
- [8] M. Li, J. Klejsa, W.B. Kleijn, On distribution preserving quantization, *IEEE Signal Process. Lett.* 17 (12) (2011) 1014–1017.
- [9] D. Simon, J. Sulam, Y. Romano, Y. Lu, M. Elad, MMSE approximation for sparse coding algorithms using stochastic resonance, *IEEE Trans. Signal Process.* 67 (17) (2019) 4597–4610.
- [10] A. Sanzi, A. Vosoughi, Distributed vector estimation for power- and bandwidth-constrained wireless sensor networks, *IEEE Trans. Signal Process.* 64 (15) (2016) 3879–3894.
- [11] H. Chen, P.K. Varshney, Nonparametric one-bit quantizers for distributed estimation, *IEEE Trans. Signal Process.* 58 (7) (2010) 3777–3787.
- [12] A. Gershman, J. Böhme, A pseudo-noise approach to direction finding, *Signal Process.* 71 (1) (1998) 1–13.
- [13] O. Dabeer, E. Masry, Multivariate signal parameter estimation under dependent noise from 1-bit dithered quantized data, *IEEE Trans. Inf. Theory* 54 (4) (2008) 1637–1654.
- [14] S. Engelberg, An introduction to dither, *IEEE Instrum. Meas. Mag.* 15 (6) (2012) 50–54.
- [15] J. Liu, F. Duan, F. Chapeau-Blondeau, D. Abbott, Exploring Bona fide optimal noise for Bayesian parameter estimation, *IEEE Access* 8 (2020) 18822–18831.
- [16] R. Benzi, A. Suter, A. Vulpiani, The mechanism of stochastic resonance, *J. Phys. A, Math. Gen.* 14 (11) (1981) L453–L457.
- [17] L. Gammaitoni, P. Hänggi, P. Jung, F. Marchesoni, Stochastic resonance, *Rev. Mod. Phys.* 70 (1) (1998) 233–287.
- [18] N.G. Stocks, Suprathreshold stochastic resonance in multilevel threshold systems, *Phys. Rev. Lett.* 84 (11) (2000) 2310–2313.

- [19] M.D. McDonnell, N.G. Stocks, C.E.M. Pearce, D. Abbott, *Stochastic Resonance: From Suprathreshold Stochastic Resonance to Stochastic Signal Quantization*, Cambridge University Press, New York, 2008.
- [20] Q. Zhai, Y. Wang, Noise effect on signal quantization in an array of binary quantizers, *Signal Process.* 152 (11) (2018) 265–272.
- [21] S. Kay, Can detectability be improved by adding noise?, *IEEE Signal Process. Lett.* 7 (1) (2000) 8–10.
- [22] H. Chen, P.K. Varshney, S.M. Kay, J.H. Michels, Theory of the stochastic resonance effect in signal detection: part I—fixed detectors, *IEEE Trans. Signal Process.* 55 (7) (2007) 3172–3184.
- [23] H. Chen, P.K. Varshney, Theory of the stochastic resonance effect in signal detection—part II: variable detectors, *IEEE Trans. Signal Process.* 56 (10) (2008) 5031–5041.
- [24] S. Lu, Q. He, F. Kong, Effects of underdamped step-varying second-order stochastic resonance for weak signal detection, *Digit. Signal Process.* 36 (1) (2015) 93–103.
- [25] F. Chapeau-Blondeau, D. Rousseau, Noise-enhanced performance for an optimal Bayesian estimator, *IEEE Trans. Signal Process.* 52 (5) (2004) 1327–1334.
- [26] H. Chen, P.K. Varshney, J.H. Michels, Noise enhanced parameter estimation, *IEEE Trans. Signal Process.* 56 (10) (2008) 5074–5081.
- [27] J. Zhu, X. Li, R.S. Blum, Y. Gu, Parameter estimation from quantized observations in multiplicative noise environments, *IEEE Trans. Signal Process.* 63 (15) (2015) 4037–4050.
- [28] A. Sani, A. Vosoughi, Noise enhanced distributed Bayesian estimation, in: *IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, New Orleans, LA, USA, 2017, pp. 4217–4221.
- [29] A. Patel, B. Kosko, Optimal mean-square noise benefits in quantizer-array linear estimation, *IEEE Signal Process. Lett.* 17 (12) (2010) 1005–1009.
- [30] T. Yang, S. Liu, H. Liu, M. Tang, X. Tan, X. Zhou, Noise benefits parameter estimation in LMMSE sense, *Digit. Signal Process.* 73 (2) (2018) 153–163.
- [31] G.O. Balkan, S. Gezici, CRLB based optimal noise enhanced parameter estimation using quantized observations, *IEEE Signal Process. Lett.* 17 (5) (2010) 477–480.
- [32] S. Uhlich, Bayes risk reduction of estimators using artificial observation noise, *IEEE Trans. Signal Process.* 63 (20) (2015) 5535–5545.
- [33] S. Ikemoto, F. Dallalibera, K. Hosoda, Noise-modulated neural networks as an application of stochastic resonance, *Neurocomputing* 277 (2) (2017) 29–37.
- [34] F. Duan, Y. Pan, F. Chapeau-Blondeau, D. Abbott, Noise benefits in combined nonlinear Bayesian estimators, *IEEE Trans. Signal Process.* 67 (17) (2019) 4611–4623.
- [35] X. Liu, L. Duan, F. Duan, F. Chapeau-Blondeau, D. Abbott, Enhancing threshold neural network via suprathreshold stochastic resonance for pattern classification, *Phys. Lett. A* 403 (2021) 127387.
- [36] A. Ribeiro, G. Giannakis, Bandwidth-constrained distributed estimation for wireless sensor networks—part I: Gaussian case, *IEEE Trans. Signal Process.* 54 (3) (2006) 1131–1143.
- [37] A. Ribeiro, G. Giannakis, Bandwidth-constrained distributed estimation for wireless sensor networks—part II: unknown probability density function, *IEEE Trans. Signal Process.* 54 (7) (2006) 2784–2796.
- [38] K. You, Recursive algorithms for parameter estimation with adaptive quantizer, *Automatica* 52 (2015) 192–201.
- [39] S.M. Kay, *Fundamentals of Statistical Signal Processing: Estimation Theory*, vol. 1, Prentice-Hall, Upper Saddle River, New Jersey, 1998.
- [40] H.L. van Trees, K.L. Bell, Z. Tan, *Detection, Estimation and Modulation Theory—Part I: Detection, Estimation, and Filtering Theory*, Wiley-Interscience, 2002.
- [41] D.K.W. Ho, B.D. Rao, Antithetic dithered 1-bit massive MIMO architecture: Efficient channel estimation via parameter expansion and PML, *IEEE Trans. Signal Process.* 67 (9) (2019) 2291–2301.
- [42] H. Chen, L.R. Varshney, P.K. Varshney, Noise-enhanced information systems, *Proc. IEEE* 102 (10) (2014) 1607–1621.
- [43] S. Bayram, S. Gezici, H.V. Poor, Noise enhanced hypothesis-testing in the restricted Bayesian framework, *IEEE Trans. Signal Process.* 58 (8) (2010) 3972–3989.
- [44] J. Nocedal, S.J. Wright, *Numerical Optimization*, Springer-Verlag, New York, 2006.
- [45] A. Sani, A. Vosoughi, On distributed linear estimation with observation model uncertainties, *IEEE Trans. Signal Process.* 66 (12) (2018) 3212–3227.

Jia Liu was born in China in 1995. She received the B. Eng. degree in communication engineering from Qingdao University in 2018 and she is currently working toward the Master degree in system science at Qingdao University. Her research interests are in signal estimation and machine learning.

Fabing Duan was born in China in 1974. He received the Master degree in engineering mechanics from the China University of Mining and Technology (Beijing) in 1999. He received, in 2002, the Ph.D. degree in solid mechanics at Zhejiang University, China. From 2002 to 2003, he was a postdoctoral fellow at the University of Angers, France. Since 2004, he has been with Qingdao University, China, and is currently a professor of system science. His research interests are in nonlinear systems and signal processing.

François Chapeau-Blondeau was born in France in 1959. He received the Engineer Diploma from ESEO, Angers, France, in 1982, the Ph.D. degree in electrical engineering from University Pierre et Marie Curie, Paris 6, France, in 1987, and the Habilitation degree from the University of Angers, France, in 1994. In 1988, he was a research associate in the Department of Biophysics at the Mayo Clinic, Rochester, Minnesota, USA, working on biomedical ultrasonics. Since 1990, he has been with the University of Angers, France, where he is currently a professor of electrical and electronic engineering. His research interests include information theory, signal processing and imaging, and the interactions between physics and information sciences.

Derek Abbott was born in South Kensington, London, U.K., in 1960. He received the B.Sc. (Hons.) degree in physics from Loughborough University, Leicestershire, U.K., in 1982 and the Ph.D. degree in electrical and electronic engineering from The University of Adelaide, Adelaide, SA, Australia, in 1995, under K. Eshraghian and B. R. Davis.

From 1978 to 1986, he was a Research Engineer with the GEC Hirst Research Centre, London, U.K. From 1986 to 1987, he was a VLSI Design Engineer with Austek Microsystems, Australia. Since 1987, he has been with The University of Adelaide, where he is currently a full Professor with the School of Electrical and Electronic Engineering. He has authored or coauthored more than 1000 publications and a number of patents. His research interests include multidisciplinary physics and electronic engineering applied to complex systems, networks, game theory, energy policy, stochastic, and biophotonics.

Prof. Abbott has been an invited speaker at more than 100 institutions. He coedited *Quantum Aspects of Life* (Imperial College Press, 2008), and coauthored *Stochastic Resonance* (Cambridge Univ. Press, 2008) and *Terahertz Imaging for Biomedical Applications* (Springer-Verlag, 2012). He is a Fellow of the Institute of Physics. He has been an Editor and/or Guest Editor for a number of journals, including the *IEEE Journal of Solid-State Circuits*, *Journal of Optics B*, *Microelectronics Journal*, *PLOS ONE*, *Proceedings of the IEEE*, and the *IEEE Photonics Journal*. He is currently on the Editorial Boards of *IEEE Access*, *Nature's Scientific Reports*, *Royal Society Open Science*, and *Frontiers in Physics*.

Prof. Abbott has received a number of awards, including the Tall Poppy Award for Science (2004), the Premier's SA Great Award in Science and Technology for outstanding contributions to South Australia (2004), an Australian Research Council Future Fellowship (2012), the David Dewhurst Medal (2015), the Barry Inglis Medal (2018), and the M. A. Sargent Medal (2019) for eminence in engineering.