

PRIOR-EXPLOITING DIRECTION-OF-ARRIVAL ALGORITHM FOR PARTIALLY UNCORRELATED SOURCE SIGNALS

P. Wirfält*, M. Jansson

KTH Royal Institute of Technology
ACCESS Linnaeus Center, Electrical
Engineering, Signal Processing
Osquidas v.10, SE-100 44 Stockholm, Sweden

G. Bouleux

University of Lyon and
University of Saint Etienne
20 Avenue de Paris
42334 Roanne Cedex, France

ABSTRACT

In certain direction-of-arrival (DOA) estimation scenarios some of the source directions are known to the operator even before measurements are acquired. It is then undesirable to use regular DOA-algorithms which waste data-samples estimating the known directions. Additionally, in some applications it is known that the signals emanating from the known directions are uncorrelated with those coming from the unknown directions. In this article we present a novel algorithm which exploits the combination of such prior knowledge in a manner more efficient (in terms of accuracy) than any algorithm known to the authors. Through numerical Monte-Carlo simulations we show the estimator to attain the theoretical accuracy bound for significantly lower signal-to-noise ratios than current state-of-the-art methods. Additionally we show the proposed algorithm to treat the stricter problem of entirely uncorrelated emitters better than current state of the art.

Index Terms— Accuracy, Arrays, Covariance matrix, Direction of arrival estimation, Signal processing algorithms

1. INTRODUCTION

In certain direction-of-arrival (DOA) estimation scenarios some of the source direction are known to the operator even before measurements are acquired. Hence there is no need to estimate such directions. Actually, in many cases the presence of such known emitters has an adverse effect on the estimation of the unknown sources [1]. Thus there is a need to remove or at least mitigate this negative effect.

This problem has recently acquired some interest; in [2] it was shown how to exploit the benefits of such prior information by modifying the sample covariance matrix of the received data when the number of available data samples were

low. A different approach is taken in [1], where asymptotically efficient (in the number of samples) methods are investigated. It was shown that very large performance gains were possible when coupling the prior bearing knowledge with the information that the sources are uncorrelated.

In this article we investigate the case when some directions are known, and it is additionally known that the signals from the emitters corresponding to the known directions are uncorrelated with the signals from the unknown sources. This assumption was implemented in a Root-MUSIC-type algorithm in [1], but it was seen that the thus defined estimator had highly undesirable properties in some scenarios. In this work we exploit the same assumption, but in a more structured way, and we derive a weighted-subspace type [3] algorithm POWDER (Prior-exploiting Orthogonally Weighted Direction Estimator), which treats the resulting problem in a statistically optimal manner.

We use the notation that T denotes transpose, c conjugate, and * conjugate-transpose. Uppercase and lowercase boldface letters denote matrices and vectors, respectively. A single index on a matrix denotes the indexed column.

2. PROBLEM DESCRIPTION

Consider the narrow-band signal model (see, e.g. [3])

$$\mathbf{y}(t) = \mathbf{A}(\bar{\boldsymbol{\theta}})\mathbf{x}(t) + \mathbf{n}(t), \quad t = 0, \dots, N-1. \quad (1)$$

The vector $\mathbf{y}(t) \in \mathbb{C}^{m \times 1}$ represents the sensor array output from a uniform linear array (ULA) with half-wavelength intra-sensor separation, and $\mathbf{x}(t) \in \mathbb{C}^{d \times 1}$ represents the signal samples. The matrix $\mathbf{A}(\bar{\boldsymbol{\theta}}) \in \mathbb{C}^{m \times d}$ is the array steering matrix, whose i th column is given by

$$\mathbf{a}(\bar{\theta}_i) = [1 \quad e^{-j\pi \sin(\bar{\theta}_i)} \quad \dots \quad e^{-j(m-1)\pi \sin(\bar{\theta}_i)}]^T, \quad (2)$$

which is uniquely determined by the array geometry and the DOAs $\bar{\boldsymbol{\theta}}$ of the impinging signals (we reserve $\boldsymbol{\theta}$ for the *unknown* DOAs, see below). The dimension m corresponds to the number of sensors and d is the number of source signals.

*Corresponding author, email: wirfalt@kth.se. The research leading to these results has received funding from the European Research Council under the European Community's Seventh Framework Programme (FP7/2007-2013) / ERC grant agreement #228044 and was partially supported by the Swedish Research Council under contract 621-2011-5847.

Finally, $\mathbf{n}(t) \in \mathbb{C}^{m \times 1}$ represents the sensor noise. We model both the signal and the noise vectors as zero mean, i.i.d. circularly symmetric complex Gaussian random processes with spatial covariance matrices given by $\text{cov}(\mathbf{x}(t)) = \mathbf{P}$ and $\text{cov}(\mathbf{n}(t)) = \sigma^2 \mathbf{I}$, respectively.

In this article we exploit two additional properties of the received data: first, we assume that some of the signal directions are known *a-priori*; hence we are only interested in estimating $d_u = d - d_k$ of the DOAs, where the subscripts u and k henceforth denotes unknown and known. With that fact in mind we can, without loss of generality, write

$$\bar{\boldsymbol{\theta}} = [\boldsymbol{\theta}^T \quad \boldsymbol{\vartheta}^T]^T; \quad (3)$$

$$\mathbf{A}(\bar{\boldsymbol{\theta}}) = [\mathbf{A}(\boldsymbol{\theta}) \quad \mathbf{A}(\boldsymbol{\vartheta})] \triangleq [\mathbf{A}_u \quad \mathbf{A}_k]; \quad (4)$$

$$\mathbf{P} = \begin{bmatrix} \mathbf{P}_u & \mathbf{P}_{uk} \\ \mathbf{P}_{uk}^* & \mathbf{P}_k \end{bmatrix} \quad (5)$$

where $\boldsymbol{\theta}$ and $\boldsymbol{\vartheta}$ denote the DOA parameters of the unknown and known sources, respectively. The second property we exploit is that we assume that there is no correlation between the signals from the known and unknown directions: hence, in (5),

$$\mathbf{P}_{uk} = \mathbf{0}. \quad (6)$$

We do not make any assumptions on $\mathbf{P}_k$ or $\mathbf{P}_u$.

3. PROPOSED ESTIMATOR: POWDER

Denote the covariance matrix of the sensor output by $\mathbf{R} = \mathbb{E}[\mathbf{y}(t)\mathbf{y}^*(t)]$. Using the notation from Section 2 we can write

$$\mathbf{R} - \sigma^2 \mathbf{I} = \mathbf{A} \mathbf{P} \mathbf{A}^* = [\mathbf{A}_u \quad \mathbf{A}_k] \begin{bmatrix} \mathbf{P}_u & \mathbf{P}_{uk} \\ \mathbf{P}_{uk}^* & \mathbf{P}_k \end{bmatrix} \begin{bmatrix} \mathbf{A}_u^* \\ \mathbf{A}_k^* \end{bmatrix}. \quad (7)$$

Let $\Pi_{\mathbf{A}_k}^\perp = \mathbf{I} - \mathbf{A}_k (\mathbf{A}_k^* \mathbf{A}_k)^{-1} \mathbf{A}_k^*$. Then $\mathbf{A}_k^* \Pi_{\mathbf{A}_k}^\perp = \mathbf{0}$, and [4], [1]

$$\begin{aligned} (\mathbf{R} - \sigma^2 \mathbf{I}) \Pi_{\mathbf{A}_k}^\perp &= \mathbf{A}_u \mathbf{P}_u \mathbf{A}_u^* \Pi_{\mathbf{A}_k}^\perp + \mathbf{A}_k \mathbf{P}_{uk}^* \mathbf{A}_u^* \Pi_{\mathbf{A}_k}^\perp \\ &= \mathbf{A}_u \mathbf{P}_u \mathbf{A}_u^* \Pi_{\mathbf{A}_k}^\perp = \mathbf{U}_s \boldsymbol{\Sigma}_s \mathbf{V}_s^*, \end{aligned} \quad (8)$$

where the second equality above is due to (6), and the third follows from a singular value decomposition where we have retained the terms corresponding to the d'_u ($\text{rank}(\mathbf{P}_u) = d'_u$) principal singular values (the remaining $m - d'_u$ singular values are all zero). The subscript s denotes signal.

Since we do not have access to the true covariance matrix $\mathbf{R}$, we instead base our estimation on the sample covariance matrix

$$\hat{\mathbf{R}} = \frac{1}{N} \sum_{t=1}^N \mathbf{y}(t) \mathbf{y}^*(t). \quad (9)$$

In order to estimate σ^2 , we perform an eigendecomposition of (9) according to

$$\hat{\mathbf{R}} = \hat{\mathbf{E}}_s \hat{\boldsymbol{\Lambda}}_s \hat{\mathbf{E}}_s^* + \hat{\mathbf{E}}_n \hat{\boldsymbol{\Lambda}}_n \hat{\mathbf{E}}_n^*, \quad (10)$$

where $\hat{\mathbf{E}}_s$ is constructed from the eigenvectors associated with the $d' = d'_u + \text{rank}(\mathbf{P}_k)$ largest eigenvalues of $\hat{\mathbf{R}}$, and $\hat{\boldsymbol{\Lambda}}_s$ contains said eigenvalues; $\hat{\mathbf{E}}_n$ spans the noise subspace, where the associated eigenvalues are $\hat{\boldsymbol{\Lambda}}_n$. Thus we can estimate $\hat{\sigma}^2 = \frac{1}{m-d'} \text{Tr}(\hat{\boldsymbol{\Lambda}}_n)$. Using this estimate and the sample data (9) in (8) and then performing the SVD we get

$$(\hat{\mathbf{R}} - \hat{\sigma}^2 \mathbf{I}) \Pi_{\mathbf{A}_k}^\perp = \hat{\mathbf{U}}_s \hat{\boldsymbol{\Sigma}}_s \hat{\mathbf{V}}_s^* + \hat{\mathbf{U}}_n \hat{\boldsymbol{\Sigma}}_n \hat{\mathbf{V}}_n^*. \quad (11)$$

The terms subscripted by n are due to noise and finite sample effects (which can be expected to be prevalent, since the sample counterpart of (6), $\hat{\mathbf{P}}_{uk} \neq \mathbf{0}$). From (11), and the orthogonality between $\hat{\mathbf{V}}_s$ and $\hat{\mathbf{V}}_n^*$,

$$\hat{\mathbf{U}}_s = (\hat{\mathbf{R}} - \hat{\sigma}^2 \mathbf{I}) \Pi_{\mathbf{A}_k}^\perp \hat{\mathbf{V}}_s \hat{\boldsymbol{\Sigma}}_s^{-1}. \quad (12)$$

By introducing the matrix $\mathbf{B}$, which spans the nullspace of $\mathbf{A}_u^*$ (i.e. $\mathbf{B}^* \mathbf{A}_u = \mathbf{0}$), we can see that

$$\begin{aligned} \mathbf{B}^* \hat{\mathbf{U}}_s &= \mathbf{B}^* (\hat{\mathbf{R}} - \hat{\sigma}^2 \mathbf{I}) \Pi_{\mathbf{A}_k}^\perp \hat{\mathbf{V}}_s \hat{\boldsymbol{\Sigma}}_s^{-1} \\ &\simeq \mathbf{B}^* (\tilde{\mathbf{R}} - \tilde{\sigma}^2 \mathbf{I}) \Pi_{\mathbf{A}_k}^\perp \mathbf{V}_s \boldsymbol{\Sigma}_s^{-1}, \end{aligned} \quad (13)$$

since $\mathbf{B}^* \mathbf{U}_s = \mathbf{0}$. In (13) $\simeq$ denotes that only first-order error terms are retained and, for a given variable, $\tilde{\mathbf{a}} = \hat{\mathbf{a}} - \mathbf{a}$ denotes the error in the estimate of that variable using the available data samples. From (2) and [5], we define $\mathbf{B} \in \mathbb{C}^{m \times m-d_u}$

$$\mathbf{B}(\boldsymbol{\theta}) = \begin{bmatrix} b_0 & b_1 & \dots & b_{d_u} & \dots & \mathbf{0} \\ & \ddots & \ddots & & \ddots & \\ \mathbf{0} & & b_0 & b_1 & \dots & b_{d_u} \end{bmatrix}^T, \quad (14)$$

where the coefficients b_i are defined by the polynomial

$$b_0 \prod_{i=1}^{d_u} (z - e^{-j\pi \sin(\theta_i)}) \triangleq b_0 z^{d_u} + b_1 z^{d_u-1} + \dots + b_{d_u}. \quad (15)$$

Since the true DOAs $\boldsymbol{\theta}$ uniquely parameterize $\mathbf{B}$ through (15), minimizing (13) with respect to $\mathbf{B}$ gives an estimate of $\boldsymbol{\theta}$. To perform this minimization in a theoretically sound manner, form the residual vector $\boldsymbol{\epsilon} = \text{vec}(\hat{\mathbf{U}}_s^* \mathbf{B})$, and solve

$$\hat{\boldsymbol{\theta}} = \arg \min_{\boldsymbol{\theta}} \boldsymbol{\epsilon}^* \mathbf{W} \boldsymbol{\epsilon}, \quad (16)$$

where $\mathbf{W}$ is a weighting matrix. In order to explicitly minimize (16) we collect the coefficients of (15) in $\mathbf{b} = [b_0 \quad b_1 \quad \dots \quad b_{d_u}]^T$ and then rewrite the residual as

$$\boldsymbol{\epsilon} = \text{vec}(\hat{\mathbf{U}}_s^* \mathbf{B}) = (\mathbf{I}_{m-d} \otimes \hat{\mathbf{U}}_s^*) \text{vec}(\mathbf{B}) \triangleq \mathbf{K} \mathbf{b}, \quad (17)$$

where $\mathbf{K} = (\mathbf{I}_{m-d} \otimes \hat{\mathbf{U}}_s^*) \boldsymbol{\Psi}$ in which the selection matrix $\boldsymbol{\Psi}$ is given from $\text{vec}(\mathbf{B}) = \boldsymbol{\Psi} \mathbf{b}$. We also use the identity $\text{vec}(\mathbf{ABC}) = (\mathbf{C}^T \otimes \mathbf{A}) \text{vec}(\mathbf{B})$ for any matrices $\mathbf{A}$, $\mathbf{B}$, $\mathbf{C}$

of compatible dimensions. With (17) we can rewrite the cost function in (16) according to

$$V(\boldsymbol{\theta}) = \boldsymbol{\epsilon}^* \mathbf{W} \boldsymbol{\epsilon} = \mathbf{b}^* \mathbf{K}^* \mathbf{W} \mathbf{K} \mathbf{b}; \quad (18)$$

the minimization of $V(\boldsymbol{\theta})$ with respect to $\boldsymbol{\theta}$ can then be solved as an eigenvalue problem for $\mathbf{b}$. For some technical details regarding this minimization, please see [6]. One benefit of parameterizing (18) by $\mathbf{b}$ is that the initialization of the weighting matrix $\mathbf{W}$ then is very straightforward, as shown below in Algorithm 1.

Choosing the weighting matrix

$$\mathbf{W} = \mathbf{E}[\boldsymbol{\epsilon} \boldsymbol{\epsilon}^*]^{-1} \quad (19)$$

is well known to (see e.g. [7]) yield minimum variance estimates of $\boldsymbol{\theta}$ in (18). In order to find a closed form expression for (19), we rewrite (17) based on (13):

$$\begin{aligned} \boldsymbol{\epsilon} &= \text{vec}(\hat{\mathbf{U}}_s^* \mathbf{B}) \\ &= (\mathbf{B}^T \otimes \boldsymbol{\Sigma}_s^{-1} \mathbf{V}_s^* \boldsymbol{\Pi}_{\mathbf{A}_k}^\perp) \text{vec}(\tilde{\mathbf{R}} - \tilde{\sigma}^2 \mathbf{I}) \triangleq \mathbf{M} \tilde{\mathbf{f}}, \end{aligned} \quad (20)$$

with $\mathbf{M}$ and $\tilde{\mathbf{f}}$ naturally defined from (20). Based on an analysis similar to the one in [8], it can be shown that $\tilde{\sigma}^2 = \frac{1}{(m-d')} \text{vec}^*(\mathbf{I}_m - \mathbf{E}_s \mathbf{E}_s^*) \text{vec}(\tilde{\mathbf{R}})$, which used in (20) gives

$$\tilde{\mathbf{f}} = \text{vec}(\tilde{\mathbf{R}}) - \frac{1}{m-d'} \text{vec}(\mathbf{I}_m) \text{vec}^*(\mathbf{I}_m - \mathbf{E}_s \mathbf{E}_s^*) \text{vec}(\tilde{\mathbf{R}}).$$

We can then write

$$\mathbf{H} = \mathbf{M} \left(\mathbf{I}_{m^2} - \frac{1}{m-d'} \text{vec}(\mathbf{I}_m) \text{vec}^*(\mathbf{I}_m - \mathbf{E}_s \mathbf{E}_s^*) \right)$$

and thus $\boldsymbol{\epsilon} = \mathbf{H} \text{vec}(\tilde{\mathbf{R}})$. Together with the well known result (from e.g. [7]) $\text{cov}(\text{vec}(\tilde{\mathbf{R}})) = N^{-1} (\mathbf{R}^T \otimes \mathbf{R})$ we find

$$\mathbf{E}[\boldsymbol{\epsilon} \boldsymbol{\epsilon}^*] = \frac{1}{N} \mathbf{H} (\mathbf{R}^T \otimes \mathbf{R}) \mathbf{H}^*, \quad (21)$$

which can be shown to be full rank. In practice the factors of (21) are based on estimates from sample-data (and $\mathbf{B}$ has to be initialized); however, as long as such estimates are consistent estimates of the true quantities, it can be shown that we can use $\hat{\mathbf{W}}$ without impairing the asymptotic properties of the estimator. We summarize the proposed method in Algorithm 1.

Remark: Note that (16) is valid for *any* array geometry; the ULA implementation is however attractive since it avoids the complex multi-dimensional optimization problem typically appearing for arbitrary array geometries.

4. NUMERICAL EXAMPLES

We confirm the efficacy of the POWDER-algorithm by numerical Monte-Carlo (MC) simulations: we generate synthetic array data according to the data-model (1); the source signal $\mathbf{x}(t)$ and the noise $\mathbf{n}(t)$ are generated as realizations of a

Algorithm 1 POWDER

```

1: Input:  $\hat{\mathbf{R}}, d, d'_u, d', \boldsymbol{\vartheta}$ ; tol; itermax
2: Find (from {Input}):  $m; d_u; \boldsymbol{\Psi}; \boldsymbol{\Pi}_{\mathbf{A}_k}^\perp; \hat{\boldsymbol{\Lambda}}_n; \hat{\mathbf{E}}_s$ 
3: Estimate:  $\hat{\sigma}^2 = (m - d')^{-1} \text{Tr}(\hat{\boldsymbol{\Lambda}}_n)$ 
4: Find:  $\hat{\mathbf{U}}_s, \hat{\boldsymbol{\Sigma}}_s, \hat{\mathbf{V}}_s^*$  from the SVD of  $(\hat{\mathbf{R}} - \hat{\sigma}^2 \mathbf{I}) \boldsymbol{\Pi}_{\mathbf{A}_k}^\perp$ 
5: Initialize:  $\mathbf{B}$  from  $\mathbf{b}^T = [1 \ 0 \ \dots \ 0]$ ;  $\mathbf{K}$ ; iter = 0
6: repeat
7:   Find:  $\hat{\mathbf{W}}^{\text{iter}}$ 
8:   Find:  $\hat{\boldsymbol{\theta}}^{\text{iter}+1}$  by minimizing (18)
9:   Find:  $\mathbf{B}^{\text{iter}+1}$  from  $\hat{\boldsymbol{\theta}}^{\text{iter}+1}$ ; update  $\hat{\mathbf{M}}^{\text{iter}+1}$ 
10:  iter  $\leftarrow$  iter + 1
11: until  $(|\hat{\boldsymbol{\theta}}^{\text{iter}} - \hat{\boldsymbol{\theta}}^{\text{iter}-1}| < \text{tol})$  OR (iter > itermax)
12: Output:  $\hat{\boldsymbol{\theta}}^{\text{iter}}$ 

```

pseudo-random process with covariance matrices $\mathbf{P}$ and $\sigma^2 \mathbf{I}$, respectively. Our performance metric is the root-mean-square error

$$\text{RMSE}_i = \sqrt{\frac{1}{L} \sum_{k=1}^L (\hat{\theta}_{i,k} - \theta_i)^2}, \quad i = 1, \dots, d_u, \quad (22)$$

where L is the number of MC-realizations.

We study the case when $\boldsymbol{\vartheta} = [12^\circ \ 20^\circ]^T$ and $\boldsymbol{\theta} = [10^\circ \ 15^\circ]^T$; to start with we let the respective signals be coherent since this is typically a difficult scenario in which many (e.g., MUSIC-type [1], [9]) estimators fail. Thus,

$$\mathbf{P}_u = \mathbf{P}_k = \begin{bmatrix} 1 & 1 \\ 1 & 1 \end{bmatrix}, \quad (23)$$

$\sigma^2 = \text{SNR}^{-1}$ and $\mathbf{P}_{uk} = \mathbf{0}$. We also let $m = 10$, and $N = 100$ samples are generated to create $\hat{\mathbf{R}}$. In Fig. 1 we vary the SNR, and compare POWDER to PLEDGE [4], which is optimal for the case when some of the signal bearings are known (but when no assumptions on the signal covariance matrix can be made). We also show CRB_P [1] (which is the theoretical performance bound when some signal bearings are known, attainable by PLEDGE) and CRB_{BD} (which is the bound exploiting the block-diagonal structure of $\mathbf{P}$ in conjunction with the knowledge of some signal bearings; this bound will be explicitly derived in the journal version of this article). For comparison, we also show the bound CRB_θ for which the known sources are absent.

We can see that the additional information $\mathbf{P}_{uk} = \mathbf{0}$ has a very large impact on the estimation accuracy; POWDER gives the same accuracy as PLEDGE for an SNR that is 25dB lower. For a given SNR, POWDER gives an error that is roughly 10 times smaller than PLEDGE. We can also see that POWDER attains the theoretical accuracy bound at a lower SNR than the other method. It is also interesting to note that CRB_{BD} is not very far from CRB_θ ; this means that the detrimental effects of the known sources have been mitigated to a large extent.

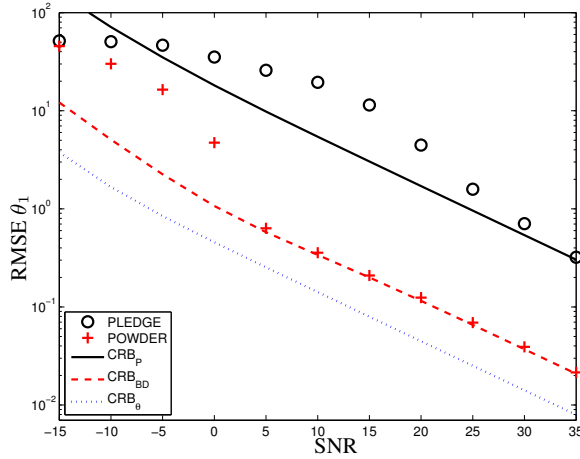


Fig. 1. Block-diagonal $\mathbf{P}$, with $\mathbf{P}_k$ and $\mathbf{P}_u$ both coherent. Two sources assumed known, $\boldsymbol{\vartheta} = [12^\circ \ 20^\circ]^T$, with $\boldsymbol{\theta} = [10^\circ \ 15^\circ]^T$. Showing RMSE of θ_1 . Averages based on 1000 MC realizations.

Next we study the case when the sources are known to be uncorrelated; we thus let $\mathbf{P} = \mathbf{I}$ and keep the remaining parameters as before, except that we increase the number of samples to $N = 1000$. In [1] it was shown that in a similar scenario a method denoted PLEDGE UC, which exploits the diagonal structure of $\mathbf{P}$ along the lines of [10], achieved large accuracy gains as compared to PLEDGE. As can be seen in Fig. 2, POWDER is asymptotically as accurate as PLEDGE UC, but attains the (joint) bound at a lower SNR. It is somewhat surprising that POWDER improves on PLEDGE UC in this scenario; the former is not exploiting as much structure in the problem as the latter is, but apparently POWDER is using its reduced information more efficiently. By comparing the (close to identical) CRBs of Fig. 2, it can be realized that the reduced parameter set (resulting from the assumption of uncorrelated sources) does not, in the studied scenario, translate to a significantly easier estimation problem as compared to the assumption of block-diagonal signal covariance. This explains the good performance of POWDER. It can also be noted that the distance to CRB_θ is larger in this scenario than in the one depicted in Fig. 1.

5. CONCLUSIONS

In this work we have derived a novel algorithm for the DOA scenario when there is prior knowledge on some directions and when it is known that the signals from the known and the unknown directions are uncorrelated. We showed through numerical simulations that the estimator, denoted POWDER, is very potent in diverse scenarios; both in the cases of co-

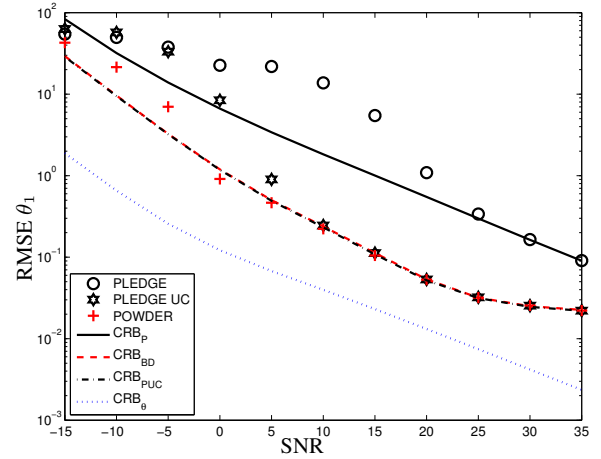


Fig. 2. Examining diagonal $\mathbf{P}$; $N = 1000$; other parameters identical to the scenario in Fig. 1.

herent and uncorrelated sources, POWDER is more accurate than any method known to the authors.

The advantages however come at a cost; the cyclic part of the algorithm, i.e. steps 6 through 11 in Algorithm 1, needs to run through more iterations than the corresponding steps of, e.g., PLEDGE or PLEDGE UC. Additionally, POWDER is not SNR-efficient in the sense that for a given N , there is a lower bound on the RMSE of the method, independent of SNR. A more thorough comparison of accuracy limitations, computational costs, and the detailed minimization operation is deferred to the journal version of this paper.

6. RELATION TO PRIOR WORK

The POWDER-method presented in this article is related to the algorithms presented in [1], [2], [9] and [11] in the sense that it exploits prior information on some source DOA. It incorporates the prior DOA-information through an orthogonal projection in a manner similar to the methods of [9] and [11] (as opposed to the methods in [1] and [2] which utilize other techniques for incorporating such knowledge). One of the methods in [1] additionally exploits information on the source correlation and the proposed method does this as well. POWDER however relaxes the assumption of the method in [1] that all sources in the scene are uncorrelated to only requiring the signals from the known directions to be uncorrelated with those from the unknown directions. It is interesting to note that even in the case (as seen in Fig. 2) when all the sources were uncorrelated, and the more restrictive assumption exploited by the method in [1] thus was satisfied, POWDER gave better performance than any previously known method.

7. REFERENCES

- [1] P. Wirfält, G. Bouleux, M. Jansson, and P. Stoica, "Optimal prior knowledge-based direction of arrival estimation," *IET Signal Process.*, vol. 6, no. 8, pp. 731–742, Oct. 2012.
- [2] J. Steinwandt, R.C. de Lamare, and M. Haardt, "Knowledge-aided direction finding based on unitary ESPRIT," in *Asilomar Conf. on Signals, Systems, and Computers*, Nov. 2011, vol. 45, pp. 613–617.
- [3] M. Viberg and B. Ottersten, "Sensor array processing based on subspace fitting," *IEEE Trans. on Signal Process.*, vol. 39, no. 5, pp. 1110–1121, May 1991.
- [4] G. Bouleux, P. Stoica, and R. Boyer, "An optimal prior knowledge-based DOA estimation method," in *17th European Signal Process. Conf.*, Glasgow, UK, Aug. 2009, pp. 869–873.
- [5] Y. Bresler and A. Macovski, "Exact maximum likelihood parameter estimation of superimposed exponential signals in noise," *IEEE Trans. on Acoust., Speech and Signal Process.*, vol. 34, no. 5, pp. 1081 – 1089, Oct. 1986.
- [6] P. Stoica and K. Sharman, "Novel eigenanalysis method for direction estimation," *IEE Proc. Part F Radar and Signal Process.*, vol. 137, no. 1, pp. 19–26, Feb. 1990.
- [7] B. Ottersten, P. Stoica, and R. Roy, "Covariance matching estimation techniques for array signal processing applications," *Digital Signal Process.*, vol. 8, no. 3, pp. 185–210, Jul. 1998.
- [8] M. Jansson, B. Göransson, and B. Ottersten, "A subspace method for direction of arrival estimation of uncorrelated emitter signals," *IEEE Trans. Signal Process.*, vol. 47, no. 4, pp. 945–956, Apr. 1999.
- [9] D.A. Linebarger, R.D. DeGroat, E.M. Dowling, P. Stoica, and G.L. Fudge, "Incorporating a priori information into MUSIC-algorithms and analysis," *Signal Process.*, vol. 46, no. 1, pp. 85–104, 1995.
- [10] M. Jansson and B. Ottersten, "Structured covariance matrix estimation: A parametric approach," in *IEEE Int. Conf. on Acoust., Speech and Signal Process.*, Istanbul, Turkey, Jun. 2000, pp. 3172–75.
- [11] R. Boyer and G. Bouleux, "Oblique projections for direction-of-arrival estimation with prior knowledge," *IEEE Trans. Signal Process.*, vol. 56, no. 4, pp. 1374–1387, Apr. 2008.