

Learning Power Spectrum Maps from Quantized Power Measurements

Daniel Romero, *Member, IEEE*, Seung-Jun Kim, *Senior Member, IEEE*,
Georgios B. Giannakis, *Fellow, IEEE*, and Roberto López-Valcarce, *Member, IEEE*

Abstract—Power spectral density (PSD) maps providing the distribution of RF power across space and frequency are constructed using power measurements collected by a network of low-cost sensors. By introducing linear compression and quantization to a small number of bits, sensor measurements can be communicated to the fusion center with minimal bandwidth requirements. Strengths of data- and model-driven approaches are combined to develop estimators capable of incorporating multiple forms of spectral and propagation prior information while fitting the rapid variations of shadow fading across space. To this end, novel nonparametric and semiparametric formulations are investigated. It is shown that PSD maps can be obtained using support vector machine-type solvers. In addition to batch approaches, an online algorithm attuned to real-time operation is developed. Numerical tests assess the performance of the novel algorithms.

I. INTRODUCTION

Power spectral density (PSD) cartography relies on sensor measurements to map the radiofrequency (RF) signal power distribution over a geographical region. The resulting maps are instrumental for various wireless network management tasks, such as power control, interference mitigation, and planning [21], [12]. For instance, PSD maps benefit wireless network planning by revealing the location of crowded regions and areas of weak coverage, thereby suggesting where new base stations should be deployed. Because they characterize how RF power distributes per channel across space, PSD maps are also useful to increase frequency reuse and mitigate interference in cellular systems. In addition, PSD maps enable opportunistic transmission in cognitive radio systems by unveiling underutilized “white spaces” in time, space, and frequency [5], [42]. Different from conventional spectrum sensing

techniques, which assume a common spectrum occupancy over the entire sensed region [28], [3], [25], PSD cartography accounts for variability across space and therefore enables a more aggressive spatial reuse.

A number of interpolation and regression techniques have been applied to construct RF power maps from power measurements. Examples include Kriging [1], compressive sampling [19], dictionary learning [23], [22], sparse Bayesian learning [18], and matrix completion [15]. These maps describe how power distributes over space but not over frequency. To accommodate variability along the frequency dimension as well, a basis expansion model was introduced in [6], [13], [8] to create PSD maps from periodograms. To alleviate the high power consumption and bandwidth needs that stem from obtaining and communicating these periodograms, [25] proposed a low-overhead sensing scheme based on single-bit data along the lines of [29]. However, this scheme assumes that the PSD is constant across space.

To summarize, existing spectrum cartography approaches either construct *power* maps from *power* measurements, or, *PSD* maps from *PSD* measurements. In contrast, the main contribution of this paper is to present algorithms capable of estimating *PSD* maps from *power* measurements, thus attaining a more efficient extraction of the information contained in the observations than existing methods. Therefore, the proposed approach enables the estimation of the RF power distribution over *frequency and space* using low-cost low-power sensors since only power measurements are required.

To facilitate practical implementations with sensor networks, where the communication bandwidth is limited, overhead is reduced by adopting two measures. First, sensor measurements are quantized to a small number of bits. Second, the available prior information is efficiently captured in both frequency and spatial domains, thus affording strong quantization while minimally sacrificing the quality of map estimates.

Specifically, a great deal of frequency-domain prior information about impinging communication waveforms can be collected from spectrum regulations and standards, which specify bandwidth, carrier frequencies, transmission masks, roll-off factors, number of sub-carriers, and so forth [38], [31]. To exploit this information, a basis expansion model is adopted, which allows the estimation of the power of each sub-channel and background noise as a byproduct. The resulting estimates can be employed to construct signal-to-noise ratio (SNR) maps, which reveal weak coverage areas, and alleviate the well-known noise uncertainty problem in cognitive radio [36].

D. Romero and G. B. Giannakis were supported by ARO grant W911NF-15-1-0492 and NSF grant 1343248. S.-J. Kim was supported by NSF grant 1547347. D. Romero and R. López-Valcarce were supported by the Spanish Ministry of Economy and Competitiveness and the European Regional Development Fund (ERDF) (projects TEC2013-47020-C2-1-R, TEC2016-76409-C2-2-R and TEC2015-69648-REDC), and by the Galician Government and ERDF (projects GRC2013/009, R2014/037 and ED431G/04).

D. Romero is with the Dept. of Information and Communication Technology, Univ. of Agder, Norway. S.-J. Kim is with the Dept. of Computer Sci. & Electrical Eng., Univ. of Maryland, Baltimore County, USA. D. Romero was and G. B. Giannakis is with the Dept. of ECE and Digital Tech. Center, Univ. of Minnesota, USA. D. Romero was and R. López-Valcarce is with the Dept. of Signal Theory and Communications, Univ. of Vigo, Spain. E-mails: daniel.romero@uia.no, sjkim@umbc.edu, georgios@umn.edu, valcarce@gts.uvigo.es.

Parts of this work have been presented at the IEEE International Conference on Acoustics, Speech, and Signal Processing, Brisbane (Australia), 2015; at the Conference on Information Sciences and Systems, Baltimore (Maryland), 2015; and at the IEEE International Workshop on Computational Advances in Multi-sensor Adaptive Processing, Cancun (Mexico), 2015.

To incorporate varying degrees of prior information in the spatial domain, nonparametric and semiparametric estimators are developed within the framework of kernel-based learning for vector-valued functions [26]. Nonparametric estimates are well suited for scenarios where no prior knowledge about the propagation environment is available due to their ability to approximate any spatial field with arbitrarily high accuracy [11]. In many cases, however, one may approximately know the transmitter locations, the path loss exponent, or even the locations of obstacles or reflectors. The proposed semiparametric estimators capture these forms of prior information through a basis expansion in the spatial domain.

Although the proposed estimators can be efficiently implemented in batch mode, limited computational resources may hinder real-time operation if an extensive set of measurements is to be processed. This issue is mitigated here through an online nonparametric estimation algorithm based on stochastic gradient descent. Remarkably, the proposed algorithm can also be applied to (vector-valued) function estimation setups besides spectrum cartography.

The present paper also contains two theoretical contributions to machine learning. First, a neat connection is established between robust (possibly vector-valued) function estimation from quantized measurements and support vector machines (SVMs) [32], [34], [35]. Through this link, theoretical guarantees and efficient implementations of SVMs permeate to the proposed methods. Second, the theory of kernel-based learning for vector-valued functions is extended to accommodate semiparametric estimation. The resulting methods are of independent interest since they can be applied beyond the present spectrum cartography context and subsume, as a special case, thin-plate splines regression [40], [8].

The rest of the paper is organized as follows. Sec. II presents the system model and formulates the problem. Sec. III proposes nonparametric and semiparametric PSD map estimation algorithms operating in batch mode, whereas Sec. IV develops an online solver. Finally, Sec. V presents some numerical tests and Sec. VI concludes the paper with closing remarks and research directions.

Notation. The cardinality of set $\mathcal{A}$ is denoted by $|\mathcal{A}|$. Scalars are denoted by lowercase letters, column vectors by boldface lowercase letters, and matrices by boldface uppercase letters. Superscript T stands for transposition, and H for conjugate transposition. The (i, j) th entry (j th column) of matrix $\mathbf{A}$ is denoted by $a_{i,j}$ ($\mathbf{a}_j$). The Kronecker product is represented by the symbol $\otimes$. The Khatri-Rao product is defined for $\mathbf{A} := [\mathbf{a}_1, \mathbf{a}_2, \dots, \mathbf{a}_N] \in \mathbb{C}^{M_1 \times N}$ and $\mathbf{B} := [\mathbf{b}_1, \mathbf{b}_2, \dots, \mathbf{b}_N] \in \mathbb{C}^{M_2 \times N}$ as $\mathbf{A} \odot \mathbf{B} := [\mathbf{a}_1 \otimes \mathbf{b}_1, \dots, \mathbf{a}_N \otimes \mathbf{b}_N] \in \mathbb{C}^{M_1 M_2 \times N}$. The entrywise or Hadamard product is defined as $(\mathbf{A} \circ \mathbf{B})_{i,j} := a_{i,j} b_{i,j}$. Vector $\mathbf{e}_{M,m}$ is the m -th column of the $M \times M$ identity matrix $\mathbf{I}_M$, whereas $\mathbf{0}_M$ and $\mathbf{1}_M$ are the vectors of dimension M with all zeros and ones, respectively. Symbol $\mathbb{E}\{\cdot\}$ denotes expectation, $\mathbb{P}\{\cdot\}$ probability, $\text{Tr}(\cdot)$ trace, $\lambda_{\max}(\cdot)$ largest eigenvalue, and $\star$ convolution. Notation $\lceil \rho \rceil$ (alternatively $\lfloor \rho \rfloor$) represents the smallest (largest) integer z satisfying $z \geq \rho$ ($z \leq \rho$).

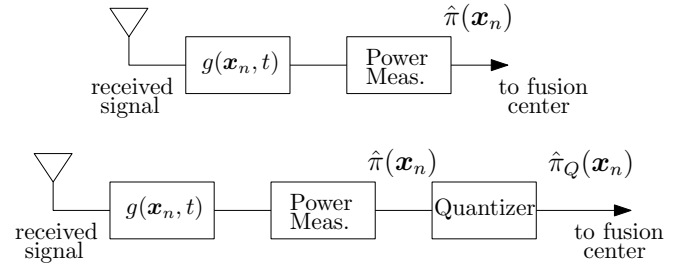


Fig. 1: Sensor architectures without (top) and with quantization (bottom).

II. SYSTEM MODEL AND PROBLEM STATEMENT

Consider $M-1$ transmitters located in a geographical region $\mathcal{R} \subset \mathbb{R}^d$, where d is typically¹ 2. Let $\phi_m(f)$ denote the transmit-PSD of the m -th transmitter and let $l_m(\mathbf{x})$ represent the gain of the channel between the m -th transmitter and location $\mathbf{x} \in \mathcal{R}$, which is assumed frequency flat for clarity; see Remark 1. If the $M-1$ transmitted signals are uncorrelated, the PSD at location $\mathbf{x}$ is given by

$$\Gamma(\mathbf{x}, f) = \sum_{m=1}^M l_m(\mathbf{x}) \phi_m(f) \quad (1)$$

where $l_M(\mathbf{x})$ is the noise power and $\phi_M(f)$ is the noise PSD, normalized to $\int_{-\infty}^{\infty} \phi_M(f) df = 1$. In view of (1), one can also normalize $\phi_m(f)$, $m = 1, \dots, M-1$, without any loss of generality to satisfy $\int_{-\infty}^{\infty} \phi_m(f) df = 1$ by absorbing any scaling factor into $l_m(\mathbf{x})$. Since such a scaling factor equals the transmit power, the coefficient $l_m(\mathbf{x})$ actually represents the power of the m -th signal at location $\mathbf{x}$.

Often in practice, the normalized PSDs $\{\phi_m(f)\}_{m=1}^{M-1}$ are approximately known since transmitters typically adhere to publicly available standards and regulations, which prescribe spectral masks, bandwidths, carrier frequencies, roll-off factors, number of pilots, and so on [38], [31]. If unknown, the methods here carry over after setting $\{\phi_m(f)\}_{m=1}^{M-1}$ to be the elements of a frequency-domain basis expansion model [6], [8]. For this reason, the rest of the paper assumes that $\{\phi_m(f)\}_{m=1}^M$ are known.

The goal is to estimate the PSD map $\Gamma(\mathbf{x}, f)$ from the measurements gathered by N sensors with locations $\{\mathbf{x}_n\}_{n=1}^N \subset \mathcal{R}$. In view of (1), this task is tantamount to estimating $\{l_m(\mathbf{x})\}_{m=1}^M$ at every spatial coordinate $\mathbf{x}$.

To minimize hardware costs and power consumption, this paper adopts the sensor architecture in Fig. 1. The *ensemble* power at the output of the filter of the n -th sensor is $\pi(\mathbf{x}_n) = \int_{-\infty}^{\infty} |G(\mathbf{x}_n, f)|^2 \Gamma(\mathbf{x}_n, f) df$, where $G(\mathbf{x}_n, f)$ denotes the frequency response of the receive filter. From (1), it follows that

$$\pi(\mathbf{x}_n) = \sum_{m=1}^M l_m(\mathbf{x}_n) \Phi_m(\mathbf{x}_n) = \boldsymbol{\phi}^T(\mathbf{x}_n) \mathbf{l}(\mathbf{x}_n) \quad (2)$$

where $\Phi_m(\mathbf{x}_n) := \int_{-\infty}^{\infty} |G(\mathbf{x}_n, f)|^2 \phi_m(f) df$ can be thought of as the contribution of the m -th transmitter per unit of

¹One may set $d = 1$ for maps along roads or railways, or even $d = 3$ for applications involving aerial vehicles or urban environments.

$l_m(\mathbf{x}_n)$ to the power of at the output of the receive filter of the n -th sensor; whereas $\phi(\mathbf{x}_n) := [\Phi_1(\mathbf{x}_n), \dots, \Phi_M(\mathbf{x}_n)]^T$ and $\mathbf{l}(\mathbf{x}_n) := [l_1(\mathbf{x}_n), \dots, l_M(\mathbf{x}_n)]^T$ are introduced to simplify notation. The n -th sensor obtains an estimate $\hat{\pi}(\mathbf{x}_n)$ of $\pi(\mathbf{x}_n)$ by measuring the signal power at the output of the filter over a certain time window. In general, $\hat{\pi}(\mathbf{x}_n) \neq \pi(\mathbf{x}_n)$ due to the finite length of this observation window.

Different from the measurement model in [6], [13], [8], where sensors obtain and communicate periodograms, the proposed scheme solely involves power measurements, thereby reducing sensor costs and bandwidth requirements. To further reduce bandwidth, $\hat{\pi}(\mathbf{x}_n)$ can be quantized as illustrated in the bottom part of Fig. 1. When uniform quantization is used, the sensors obtain

$$\hat{\pi}_Q(\mathbf{x}_n) := Q(\hat{\pi}(\mathbf{x}_n)) := \lfloor \hat{\pi}(\mathbf{x}_n) / (2\epsilon) \rfloor, \quad n = 1, \dots, N \quad (3)$$

where 2ϵ is the quantization step; see also Remark 5. If R denotes the number of quantization levels, which depends on ϵ and the range of $\hat{\pi}(\mathbf{x}_n)$, the number of bits needed is now just $\lceil \log_2 R \rceil$. Depending on how accurate $\hat{\pi}(\mathbf{x}_n)$ is, either $Q(\pi(\mathbf{x}_n)) = Q(\hat{\pi}(\mathbf{x}_n))$ or $Q(\pi(\mathbf{x}_n)) \neq Q(\hat{\pi}(\mathbf{x}_n))$. The latter event is termed *measurement error* and is due to the finite length of the aforementioned time window.

Finally, the sensors communicate the measurements $\{\hat{\pi}(\mathbf{x}_n)\}_{n=1}^N$ or $\{\hat{\pi}_Q(\mathbf{x}_n)\}_{n=1}^N$ to the fusion center. Given these measurements, together with $\{\phi_m(f)\}_{m=1}^M$ and $\{\mathbf{x}_n\}_{n=1}^N$, the problem is to estimate $\{l_m(\mathbf{x})\}_{m=1}^M$ at every $\mathbf{x} \in \mathcal{R}$. The latter can be viewed individually as M functions of the spatial coordinate $\mathbf{x}$, or, altogether as a vector field $\mathbf{l} : \mathbb{R}^d \rightarrow \mathbb{R}^M$, where $\mathbf{l}(\mathbf{x}) := [l_1(\mathbf{x}), \dots, l_M(\mathbf{x})]^T$. Thus, estimating the PSD map in (1) is in fact a problem of estimating a vector-valued function.

Remark 1 (Frequency-selective channels). *If the channels are not frequency-flat, then each term in the sum of (1) can be decomposed into multiple components of smaller bandwidth in such a way that each one sees an approximately frequency-flat channel. To ensure that these components are uncorrelated, one can choose their frequency supports to be disjoint.*

Remark 2. *Sensors can also operate digitally. In this case, one could implement the receive filters to have pseudo-random impulse responses [25]. Selecting distinct seeds for the random number generators of different sensors yields linearly independent $\{\phi(\mathbf{x}_n)\}_{n=1}^N$ with high probability, which ensures identifiability of $\{l_m(\mathbf{x}_n)\}_{m=1}^M$ (cf. (2)).*

Remark 3. *If a wideband map is to be constructed, then Nyquist-rate sampling may be too demanding for low-cost sensors. In this scenario, one can replace the filter in Fig. 1 with an analog-to-information-converter (A2IC) [37], [31]. To see that (2) still holds and therefore the proposed schemes still apply, let $\mathbf{G}(\mathbf{x}_n)$ represent the compression matrix of the n th sensor, which multiplies raw measurement blocks to yield compressed data blocks [31]. The ensemble power of the latter is proportional to $\pi(\mathbf{x}_n) := \text{Tr}(\mathbf{G}(\mathbf{x}_n)\Sigma(\mathbf{x}_n)\mathbf{G}^T(\mathbf{x}_n))$, where $\Sigma(\mathbf{x}_n) = \sum_{m=1}^M l_m(\mathbf{x}_n)\Sigma_m$ denotes the covariance*

matrix of the uncompressed data blocks, and Σ_m the covariance matrix of the blocks transmitted from the m -th transmitter. Combining both equalities yields (2) upon defining $\Phi_m(\mathbf{x}_n) := \text{Tr}(\mathbf{G}(\mathbf{x}_n)\Sigma_m\mathbf{G}^T(\mathbf{x}_n))$.

III. LEARNING PSD MAPS

This section develops various PSD map estimators offering different bandwidth-performance trade-offs. First, Sec. III-A puts forth a nonparametric estimator to recover $\mathbf{l}(\mathbf{x}) := [l_1(\mathbf{x}), \dots, l_M(\mathbf{x})]^T$ from un-quantized measurements. To reduce bandwidth requirements, this approach is extended in Sec. III-B to accommodate quantized data. The detrimental impact of strong quantization on the quality of map estimates is counteracted in Sec. III-C by leveraging propagation prior information. For simplicity, these methods are presented for the scenario where each sensor obtains a single measurement, whereas general versions accommodating multiple measurements per sensor are outlined in Sec. III-E.

A. Estimation via nonparametric regression

This section reviews the background in kernel-based learning necessary to develop the cartography tools in the rest of the paper and presents an estimator to obtain PSD maps from un-quantized measurements.

Kernel-based regression seeks estimates among a wide class of functions termed *reproducing kernel Hilbert space* (RKHS). In the present setting of vector-valued functions, such an RKHS is given by $\mathcal{H} := \{\mathbf{l}(\mathbf{x}) = \sum_{n=1}^{\infty} \mathbf{K}(\mathbf{x}, \tilde{\mathbf{x}}_n) \tilde{\mathbf{c}}_n : \tilde{\mathbf{x}}_n \in \mathbb{R}^d, \tilde{\mathbf{c}}_n \in \mathbb{R}^M\}$ [26], [11], where $\{\tilde{\mathbf{c}}_n\}_{n=1}^{\infty}$ are $M \times 1$ expansion coefficient vectors and $\mathbf{K}(\mathbf{x}, \mathbf{x}')$ is termed reproducing kernel map. The latter is any matrix-valued function $\mathbf{K}(\mathbf{x}, \mathbf{x}') : \mathbb{R}^d \times \mathbb{R}^d \rightarrow \mathbb{R}^{M \times M}$ that is (i) symmetric, meaning that $\mathbf{K}(\mathbf{x}, \mathbf{x}') = \mathbf{K}(\mathbf{x}', \mathbf{x})$ for any $\mathbf{x}$ and $\mathbf{x}'$; and (ii) positive (semi)definite, meaning that the square matrix having $\mathbf{K}(\tilde{\mathbf{x}}_n, \tilde{\mathbf{x}}_{n'})$ as its (n, n') block is positive semi-definite for any $\{\tilde{\mathbf{x}}_1, \dots, \tilde{\mathbf{x}}_{\tilde{N}}\}$. Remark 4 guides the selection of functions $\mathbf{K}(\mathbf{x}, \mathbf{x}')$ qualifying as reproducing kernels.

As any Hilbert space, an RKHS has an associated inner product, not necessarily equal to the classical $\langle \mathbf{f}, \mathbf{g} \rangle = \int \mathbf{f}^T(\mathbf{x})\mathbf{g}(\mathbf{x})d\mathbf{x}$. Specifically, the inner product between two RKHS functions $\mathbf{l}(\mathbf{x}) := \sum_{n=1}^{\tilde{N}} \mathbf{K}(\mathbf{x}, \tilde{\mathbf{x}}_n) \tilde{\mathbf{c}}_n$ and $\mathbf{l}'(\mathbf{x}) := \sum_{n=1}^{\tilde{N}'} \mathbf{K}(\mathbf{x}, \tilde{\mathbf{x}}'_n) \tilde{\mathbf{c}}'_n$ can be obtained through the reproducing kernel as

$$\langle \mathbf{l}, \mathbf{l}' \rangle_{\mathcal{H}} = \sum_{n=1}^{\tilde{N}} \sum_{n'=1}^{\tilde{N}'} \tilde{\mathbf{c}}_n^T \mathbf{K}(\tilde{\mathbf{x}}_n, \tilde{\mathbf{x}}'_{n'}) \tilde{\mathbf{c}}'_{n'}. \quad (4)$$

This expression is referred to as the *reproducing property* and is of paramount importance since it allows the computation of function inner products without integration. From (4), the induced RKHS norm of $\mathbf{l}$ can be written as $\|\mathbf{l}\|_{\mathcal{H}}^2 := \langle \mathbf{l}, \mathbf{l} \rangle = \sum_{n=1}^{\tilde{N}} \sum_{n'=1}^{\tilde{N}} \tilde{\mathbf{c}}_n^T \mathbf{K}(\tilde{\mathbf{x}}_n, \tilde{\mathbf{x}}_{n'}) \tilde{\mathbf{c}}_{n'}$ and is widely used as a proxy for smoothness of $\mathbf{l}$.

Kernel-based methods confine their search of estimates to functions in $\mathcal{H}$, which is not a limitation since an extensive class of functions, including any continuous function vanishing at infinity, can be approximated with arbitrary accuracy by a

function in $\mathcal{H}$ for a properly selected kernel [11]. However, function estimation is challenging since (i) any finite set of samples generally admits infinitely many interpolating functions in $\mathcal{H}$ and (ii) the estimate $\hat{\mathbf{l}}$ does not generally approach the estimated function $\mathbf{l}$ even for an arbitrary large number of noisy samples if the latter are overfitted. To mitigate both issues, kernel regression seeks estimates minimizing the sum of two terms, where the first penalizes estimates deviating from the observations and the second promotes smoothness. Specifically, if $\mathcal{L}(e_n)$ denotes a loss function of the measurement error $e_n := \hat{\pi}(\mathbf{x}_n) - \phi^T(\mathbf{x}_n)\mathbf{l}(\mathbf{x}_n)$, the nonparametric kernel-based regression estimate of $\mathbf{l}$ is [9]

$$\hat{\mathbf{l}} := \arg \min_{\mathbf{l} \in \mathcal{H}} \frac{1}{N} \sum_{n=1}^N \mathcal{L}(\hat{\pi}(\mathbf{x}_n) - \phi^T(\mathbf{x}_n)\mathbf{l}(\mathbf{x}_n)) + \lambda \|\mathbf{l}\|_{\mathcal{H}}^2 \quad (5)$$

where the user-selected scalar $\lambda > 0$ controls the tradeoff between fitting the data and smoothness of the estimate, captured by its RKHS norm.

Since $\mathcal{H}$ is infinite-dimensional, solving (5) directly is generally not possible with a finite number of operations. Fortunately, the so-called *representer theorem* (see e.g., [32], [2]) asserts that $\hat{\mathbf{l}}$ in (5) is of the form

$$\sum_{n=1}^N \mathbf{K}(\mathbf{x}, \mathbf{x}_n) \mathbf{c}_n := \mathbf{K}(\mathbf{x}) \mathbf{c} \quad (6)$$

for some $\{\mathbf{c}_n\}_{n=1}^N$, where $\mathbf{K}(\mathbf{x}) := [\mathbf{K}(\mathbf{x}, \mathbf{x}_1), \dots, \mathbf{K}(\mathbf{x}, \mathbf{x}_N)]$ is of size $M \times MN$, and $\mathbf{c} := [\mathbf{c}_1^T, \dots, \mathbf{c}_N^T]^T$ is $MN \times 1$. In words, the solution to (5) admits a kernel expansion around the sensor locations $\{\mathbf{x}_n\}_{n=1}^N$. From (6), it follows that finding $\hat{\mathbf{l}}$ amounts to finding $\mathbf{c}$, but the latter can easily be obtained by solving the problem that results from substituting (6) into (5):

$$\hat{\mathbf{c}} := \arg \min_{\mathbf{c}} \sum_{n=1}^N \mathcal{L}(\hat{\pi}(\mathbf{x}_n) - \phi^T(\mathbf{x}_n) \mathbf{K}(\mathbf{x}_n) \mathbf{c}) + \lambda N \mathbf{c}^T \mathbf{K} \mathbf{c}. \quad (7)$$

In (7), the $MN \times MN$ matrix $\mathbf{K}$ is formed to have $\mathbf{K}(\mathbf{x}_n, \mathbf{x}_{n'})$ as its (n, n') -th block. Clearly, the minimizer of (5) can then be recovered as $\hat{\mathbf{l}}(\mathbf{x}) = \mathbf{K}(\mathbf{x}) \hat{\mathbf{c}}$.

The loss function $\mathcal{L}$ is typically chosen to be convex. The simplest example is the squared loss $\mathcal{L}_2(e_n) := e_n^2$, with the resulting $\hat{\mathbf{l}}$ referred to as the *kernel ridge regression* estimate. Defining $\hat{\pi} := [\hat{\pi}(\mathbf{x}_1), \dots, \hat{\pi}(\mathbf{x}_N)]^T$, $\Phi := [\phi(\mathbf{x}_1), \dots, \phi(\mathbf{x}_N)]$, and $\Phi_0 := \mathbf{I}_N \odot \Phi$, expression (7) becomes

$$\begin{aligned} \hat{\mathbf{c}} &= \arg \min_{\mathbf{c} \in \mathbb{R}^{MN}} \|\hat{\pi} - \Phi_0^T \mathbf{K} \mathbf{c}\|_2^2 + \lambda N \mathbf{c}^T \mathbf{K} \mathbf{c} \\ &= (\Phi_0 \Phi_0^T \mathbf{K} + \lambda N \mathbf{I}_{MN})^{-1} \Phi_0 \hat{\pi}. \end{aligned} \quad (8)$$

Besides its simplicity of implementation, the estimate $\hat{\mathbf{l}}(\mathbf{x}) = \mathbf{K}(\mathbf{x}) \hat{\mathbf{c}}$ with $\hat{\mathbf{c}}$ as in (8) offers a twofold advantage over existing cartography schemes. First, existing estimators relying on *power measurements* can only construct *power maps* [1], [19], [18], [23], [22], whereas the proposed method is capable of obtaining *PSD maps* from the same measurements. On the other hand, existing methods for estimating *PSD maps* require *PSD measurements*, that is, every sensor must obtain and

transmit periodograms to the fusion center. This necessitates a higher communication bandwidth, longer sensing time, and more costly sensors than required here [6], [13], [8].

Remark 4. *The choice of the kernel considerably affects the estimation performance when the number of observations is small. Thus, it is important to choose a kernel that is well-suited to the spatial variability of the true $\mathbf{l}$. To do so, one may rely on cross validation, historical data [32, Sec. 2.3], or multi-kernel approaches [7]. Although specifying matrix-valued kernels is more challenging than specifying their scalar counterparts ($M = 1$) [26], a simple but effective possibility is to construct a diagonal kernel as $\mathbf{K}(\mathbf{x}, \mathbf{x}') = \text{diag}(k_1(\mathbf{x}, \mathbf{x}'), \dots, k_M(\mathbf{x}, \mathbf{x}'))$ where $\{k_m(\mathbf{x}, \mathbf{x}')\}_{m=1}^M$ are valid scalar kernels. For example, $k_m(\mathbf{x}, \mathbf{x}')$ can be the popular Gaussian kernel $k_m(\mathbf{x}, \mathbf{x}') = \exp(-\|\mathbf{x} - \mathbf{x}'\|^2 / \sigma_m^2)$.*

B. Nonparametric regression from quantized data

The scheme in Sec. III-A offers a drastic bandwidth reduction relative to competing PSD map estimators since only scalar-valued measurements need to be communicated to the fusion center. The methods in this section accomplish a further reduction by accommodating quantized measurements.

Recall from Sec. II, that $\{\hat{\pi}_Q(\mathbf{x}_n)\}_{n=1}^N$ denote the result of uniformly quantizing the power measurements $\{\hat{\pi}(\mathbf{x}_n)\}_{n=1}^N$; see also Remark 5. The former essentially convey *interval information* about the latter, since (3) implies that $\hat{\pi}(\mathbf{x}_n)$ is contained in the interval $[y(\mathbf{x}_n) - \epsilon, y(\mathbf{x}_n) + \epsilon)$, where $y(\mathbf{x}_n) := [2\hat{\pi}_Q(\mathbf{x}_n) + 1]\epsilon$. Note that $y(\mathbf{x}_n)$ is in fact the centroid of the $\hat{\pi}_Q(\mathbf{x}_n)$ -th quantization interval.

To account for the uncertainty within such an interval, one can replace $\hat{\pi}(\mathbf{x}_n)$ in (5) with $y(\mathbf{x}_n)$ as

$$\hat{\mathbf{l}} = \arg \min_{\mathbf{l} \in \mathcal{H}} \frac{1}{N} \sum_{n=1}^N \mathcal{L}(y(\mathbf{x}_n) - \phi^T(\mathbf{x}_n)\mathbf{l}(\mathbf{x}_n)) + \lambda \|\mathbf{l}\|_{\mathcal{H}}^2 \quad (9)$$

and set $\mathcal{L}$ to assign no cost across all candidate functions $\mathbf{l}$ that lead to values of $\phi^T(\mathbf{x}_n)\mathbf{l}(\mathbf{x}_n) = \pi(\mathbf{x}_n)$ falling $\pm\epsilon$ around $y(\mathbf{x}_n)$. In other words, such an $\mathcal{L}$ only penalizes functions $\mathbf{l}$ for which $e_n = y(\mathbf{x}_n) - \phi^T(\mathbf{x}_n)\mathbf{l}(\mathbf{x}_n)$ falls outside of $[-\epsilon, \epsilon)$. Examples of these ϵ -insensitive loss functions include $\mathcal{L}_{1\epsilon}(e_n) := \max(0, |e_n| - \epsilon)$ [32], [34], and the less known $\mathcal{L}_{2\epsilon}(e_n) := \max(0, e_n^2 - \epsilon)$. Incidentally, these functions endow the proposed estimators with robustness to outliers and promote sparsity in $\{e_n\}_{n=1}^N$, which is being a particularly well-motivated property when the number of measurement errors is small relative to N , that is, when $Q(\pi(\mathbf{x}_n)) = Q(\hat{\pi}(\mathbf{x}_n))$ for most values of n .

The rest of this section develops solvers for (9) and establishes a link between (9) and SVMs. To this end, note that application of the representer theorem to (9) yields, as in Sec. III-A, an estimate $\hat{\mathbf{l}}(\mathbf{x}) = \mathbf{K}(\mathbf{x}) \hat{\mathbf{c}}$ with

$$\hat{\mathbf{c}} = \arg \min_{\mathbf{c}} \sum_{n=1}^N \mathcal{L}(y(\mathbf{x}_n) - \phi^T(\mathbf{x}_n) \mathbf{K}(\mathbf{x}_n) \mathbf{c}) + \lambda N \mathbf{c}^T \mathbf{K} \mathbf{c}. \quad (10)$$

Now focus on $\mathcal{L}_{1\epsilon}$ and note that $\mathcal{L}_{1\epsilon}(e_n) = \xi_n + \zeta_n$, where $\xi_n := \max(0, e_n - \epsilon)$ and $\zeta_n := \max(0, -e_n - \epsilon)$ respectively

quantify positive deviations of e_n with respect to the right and left endpoints of $[-\epsilon, \epsilon]$. This implies that ξ_n satisfies $\xi_n \geq e_n - \epsilon$ and $\xi_n \geq 0$, whereas ζ_n satisfies $\zeta_n \geq -e_n - \epsilon$ and $\zeta_n \geq 0$, thus establishing the following result.

Proposition 1. *The problem in (10) with $\mathcal{L}$ the ϵ -insensitive loss function $\mathcal{L}_{1\epsilon}$ can be expressed as*

$$\begin{aligned} (\hat{c}, \hat{\xi}, \hat{\zeta}) &= \arg \min_{c, \xi, \zeta} \sum_{n=1}^N (\xi_n + \zeta_n) + \lambda N c^T K c \quad (11) \\ \text{s.to } \xi_n &\geq y(\mathbf{x}_n) - \phi^T(\mathbf{x}_n) K(\mathbf{x}_n) c - \epsilon, \quad \xi_n \geq 0, \\ \zeta_n &\geq -y(\mathbf{x}_n) + \phi^T(\mathbf{x}_n) K(\mathbf{x}_n) c - \epsilon, \quad \zeta_n \geq 0, \\ n &= 1, \dots, N \end{aligned}$$

where $\xi := [\xi_1, \dots, \xi_N]^T$ and $\zeta := [\zeta_1, \dots, \zeta_N]^T$.

Problem (11) is a convex quadratic program with slack variables $\{\xi_n, \zeta_n\}_{n=1}^N$. Although one can obtain $(\hat{c}, \hat{\xi}, \hat{\zeta})$ using, for example, an off-the-shelf interior-point solver, it will be shown that a more efficient approach is to solve the dual-domain version of (11).

To the best of our knowledge, (11) constitutes the first application of an ϵ -insensitive loss to estimating functions from quantized data. As expected from the choice of loss function and regularizer, (11) is an SVM-type problem. However, different from existing SVMs, for which data comprises vector-valued noisy versions of $\{\mathbf{l}(\mathbf{x}_n)\}_{n=1}^N$ [26, Examples 1 and 2], the estimate in (11) relies on noisy versions of the scalars $\{\phi^T(\mathbf{x}_n) \mathbf{l}(\mathbf{x}_n)\}_{n=1}^N$. Therefore, (11) constitutes a new class of SVM for vector-valued function estimation. As a desirable consequence of this connection, the proposed estimate inherits the well-documented generalization performance of existing SVMs [26], [35], [11]. However, it is prudent to highlight one additional notable difference between (11) and conventional SVMs that pertains to the present context of function estimation from quantized data: whereas in the present setting ϵ is determined by the quantization interval length, this parameter must be delicately tuned in conventional SVMs to control generalization performance.

The proposed estimator is *nonparametric* since the number of unknowns in (11) depends on the number of observations N . Although this number of unknowns also grows with M , it is shown next that this is not the case in the dual formulation. To see this, let $K_0 := \Phi_0^T K \Phi_0$ as well as $\mathbf{y} := [y(\mathbf{x}_1), \dots, y(\mathbf{x}_N)]^T$. With α and β representing the Lagrange multipliers associated with the $\{\xi_n\}$ and the $\{\zeta_n\}$ constraints, the dual of (11) can be easily shown to be

$$\begin{aligned} (\hat{\alpha}, \hat{\beta}) &:= \arg \min_{\alpha, \beta \in \mathbb{R}^N} \frac{1}{4N\lambda} (\alpha - \beta)^T K_0 (\alpha - \beta) \\ &\quad - (\mathbf{y} - \epsilon \mathbf{1}_N)^T \alpha + (\mathbf{y} + \epsilon \mathbf{1}_N)^T \beta \\ \text{s.to } \mathbf{0}_N &\leq \alpha \leq \mathbf{1}_N, \mathbf{0}_N \leq \beta \leq \mathbf{1}_N. \quad (12) \end{aligned}$$

From the Karush-Kuhn-Tucker (KKT) conditions, the primal variables can be recovered from the dual ones using

$$\hat{c} = \frac{1}{2\lambda N} \Phi_0 (\hat{\alpha} - \hat{\beta}) \quad (13a)$$

$$\hat{\xi} = \max(\mathbf{0}_N, \mathbf{y} - \Phi_0^T K \hat{c} - \epsilon \mathbf{1}_N) \quad (13b)$$

Algorithm 1: Nonparametric batch PSD map estimator

- 1: **Input:** $\{(\mathbf{x}_n, \phi(\mathbf{x}_n), \hat{\pi}_Q(\mathbf{x}_n))\}_{n=1}^N, \{\phi_m(f)\}_{m=1}^M, \epsilon$
- 2: **Parameters:** $\lambda, K(\mathbf{x}, \mathbf{x}')$
- 3: $\Phi = [\phi(\mathbf{x}_1), \dots, \phi(\mathbf{x}_N)]$
- 4: $\Phi_0 = \mathbf{I}_N \odot \Phi$
- 5: Form K , whose (n, n') -th block is $K(\mathbf{x}_n, \mathbf{x}_{n'})$
- 6: $K_0 = \Phi_0^T K \Phi_0$
- 7: $\mathbf{y} = [y(\mathbf{x}_1), \dots, y(\mathbf{x}_N)]^T$
- 8: Obtain $(\hat{\alpha}, \hat{\beta})$ from (12)
- 9: $[\hat{c}_1^T, \dots, \hat{c}_N^T]^T = [1/(2\lambda N)] \Phi_0 (\hat{\alpha} - \hat{\beta})$
- 10: **Output:** Function $\hat{\Gamma}(\mathbf{x}, f) = \sum_{m=1}^M \hat{l}_m(\mathbf{x}) \phi_m(f)$, where $[\hat{l}_1(\mathbf{x}), \dots, \hat{l}_M(\mathbf{x})]^T = \sum_{n=1}^N K(\mathbf{x}, \mathbf{x}_n) \hat{c}_n$.

$$\hat{\zeta} = \max(\mathbf{0}_N, -\mathbf{y} + \Phi_0^T K \hat{c} - \epsilon \mathbf{1}_N). \quad (13c)$$

Algorithm 1 lists the steps involved in the proposed nonparametric estimator. Note that the primal (11) entails $(M+2)N$ variables whereas the dual (12) has just $2N$. The latter can be solved using sequential minimal optimization algorithms [27], which here can afford simplified implementation along the lines of e.g., [20] because there is no bias term. However, for moderate problem sizes ($< 5,000$), interior point solvers are more reliable [32, Ch. 10] while having *worst-case* complexity $\mathcal{O}(N^{3.5})$. As a desirable byproduct, interior point methods directly provide the Lagrange multipliers, which are useful for recovering the primal variables (cf. (23)).

C. Semiparametric regression using quantized data

The nonparametric estimators in Secs. III-A and III-B are universal in the sense that they can approximate wide classes of functions, including all continuous functions $\mathbf{l}$ vanishing at infinity, with arbitrary accuracy provided that the number of measurements is large enough [11]. However, since measurements are limited in number and can furthermore be quantized, incorporating available prior knowledge is crucial to improve the accuracy of the estimates. One could therefore consider applying *parametric* approaches since they can readily incorporate various forms of prior information. However, these approaches lack the flexibility of nonparametric techniques since they can only estimate functions in very limited classes. Semiparametric alternatives offer a “sweet spot” by combining the merits of both approaches [32].

This section presents semiparametric estimators capable of capturing prior information about the propagation environment yet preserving the flexibility of the nonparametric estimators in Secs. III-A and III-B. To this end, an estimate of the form $\mathbf{l}(\mathbf{x}) = \mathbf{l}'(\mathbf{x}) + \check{\mathbf{l}}(\mathbf{x})$ is pursued, where (cf. Secs. III-A and III-B) the nonparametric component $\mathbf{l}'(\mathbf{x})$ belongs to an RKHS $\mathcal{H}'$ with kernel matrix K' ; whereas the parametric component is given by

$$\check{\mathbf{l}}(\mathbf{x}) = \sum_{\nu=1}^{N_B} B_\nu(\mathbf{x}) \theta_\nu := B(\mathbf{x}) \boldsymbol{\theta} \quad (14)$$

with $B(\mathbf{x}) := [B_1(\mathbf{x}), \dots, B_{N_B}(\mathbf{x})]$ collecting N_B user-selected basis matrix functions $B_\nu(\mathbf{x}) : \mathcal{R} \rightarrow \mathbb{R}^{M \times M}$, $\nu = 1, \dots, N_B$, and $\boldsymbol{\theta} := [\theta_1^T, \dots, \theta_{N_B}^T]^T$.

If the transmitter locations $\{\chi_m\}_{m=1}^M$ are approximately known, the free space propagation loss can be described by matrix basis functions of the form $B_m(\mathbf{x}) = f_m(\|\mathbf{x} - \chi_m\|) \mathbf{e}_{M,m} \mathbf{e}_{M,m}^T$, where $f_m(\|\mathbf{x} - \chi_m\|)$ is the attenuation between a transmitter located at χ_m and a receiver located at an arbitrary point $\mathbf{x}$; see Sec. V for an example. Note that if this basis accurately captures the propagation effects in $\mathcal{R}$, then the m -th entry of the estimated θ_m is approximately proportional to the transmit power of the m -th transmitter.

An immediate two-step approach to estimating $\mathbf{l}$ is to first fit the data with $\tilde{\mathbf{l}}$ in (14), and then fit the residuals with $\mathbf{l}'$ as detailed in Sec. III-B. Since this so-termed *back-fitting* approach is known to yield sub-optimal estimates [32], this paper pursues a joint fit, which constitutes a novel approach in kernel-based learning for vector-valued functions. To this end, define $\mathcal{H}$ as the space of functions $\mathbf{l}$ (not necessarily an RKHS) that can be written as $\mathbf{l} = \mathbf{l}' + \tilde{\mathbf{l}}$, with $\mathbf{l}' \in \mathcal{H}'$ and $\tilde{\mathbf{l}}$ as in (14). One can thereby seek semiparametric estimates of the form

$$\hat{\mathbf{l}} = \arg \min_{\mathbf{l} \in \mathcal{H}} \frac{1}{N} \sum_{n=1}^N \mathcal{L}(y(\mathbf{x}_n) - \phi^T(\mathbf{x}_n) \mathbf{l}(\mathbf{x}_n)) + \lambda \|\mathbf{l}'\|_{\mathcal{H}'}^2 \quad (15)$$

where the regularizer involves only the nonparametric component through the norm $\|\cdot\|_{\mathcal{H}'}$ in $\mathcal{H}'$. Using [2, Th. 3.1], one can readily generalize the representer theorem in [32, Th. 4.3] to the present semiparametric case. This yields $\hat{\mathbf{l}}(\mathbf{x}) = \mathbf{K}'(\mathbf{x}) \hat{\mathbf{c}}' + \mathbf{B}(\mathbf{x}) \hat{\boldsymbol{\theta}}$, where

$$(\hat{\mathbf{c}}', \hat{\boldsymbol{\theta}}) = \arg \min_{\mathbf{c}', \boldsymbol{\theta}} \sum_{n=1}^N \mathcal{L}(y(\mathbf{x}_n) - \phi^T(\mathbf{x}_n) [\mathbf{K}'(\mathbf{x}_n) \mathbf{c}' + \mathbf{B}(\mathbf{x}_n) \boldsymbol{\theta}]) + \lambda N \mathbf{c}'^T \mathbf{K}' \mathbf{c}' \quad (16)$$

Comparing (10) with (16), and replacing $\mathbf{K}(\mathbf{x}_n) \mathbf{c}$ with $\mathbf{K}'(\mathbf{x}_n) \mathbf{c}' + \mathbf{B}(\mathbf{x}_n) \boldsymbol{\theta}$ yields the next result (cf. Proposition 1).

Proposition 2. *The problem in (16) with $\mathcal{L}$ the ϵ -insensitive loss function $\mathcal{L}_{1\epsilon}$ can be expressed as*

$$(\hat{\mathbf{c}}', \hat{\boldsymbol{\theta}}, \hat{\boldsymbol{\xi}}, \hat{\boldsymbol{\zeta}}) = \arg \min_{\mathbf{c}', \boldsymbol{\theta}, \boldsymbol{\xi}, \boldsymbol{\zeta}} \sum_{n=1}^N (\xi_n + \zeta_n) + \lambda N \mathbf{c}'^T \mathbf{K}' \mathbf{c}'$$

s.t. $\xi_n \geq y(\mathbf{x}_n) - \phi^T(\mathbf{x}_n) [\mathbf{K}'(\mathbf{x}_n) \mathbf{c}' + \mathbf{B}(\mathbf{x}_n) \boldsymbol{\theta}] - \epsilon$, $\xi_n \geq 0$

$$\zeta_n \geq -y(\mathbf{x}_n) + \phi^T(\mathbf{x}_n) [\mathbf{K}'(\mathbf{x}_n) \mathbf{c}' + \mathbf{B}(\mathbf{x}_n) \boldsymbol{\theta}] - \epsilon, \zeta_n \geq 0$$

$$n = 1, \dots, N. \quad (17)$$

The primal problem in (17) entails vectors of size MN , which motivates solving its dual version. Upon defining $\mathbf{B}$ as an $NM \times N_B M$ matrix whose (n, ν) -th block is $B_\nu(\mathbf{x}_n)$, and representing the Lagrange multipliers associated with the $\{\xi_n\}$ and $\{\zeta_n\}$ constraints by $\boldsymbol{\alpha}$ and $\boldsymbol{\beta}$, the dual of (17) is

$$(\hat{\boldsymbol{\alpha}}, \hat{\boldsymbol{\beta}}) = \arg \min_{\boldsymbol{\alpha}, \boldsymbol{\beta}} \frac{1}{4N\lambda} (\boldsymbol{\alpha} - \boldsymbol{\beta})^T \mathbf{K}'_0 (\boldsymbol{\alpha} - \boldsymbol{\beta})$$

$$- (\mathbf{y} - \epsilon \mathbf{1}_N)^T \boldsymbol{\alpha} + (\mathbf{y} + \epsilon \mathbf{1}_N)^T \boldsymbol{\beta} \quad (18)$$

s. to $\mathbf{0}_N \leq \boldsymbol{\alpha} \leq \mathbf{1}_N$, $\mathbf{0}_N \leq \boldsymbol{\beta} \leq \mathbf{1}_N$

$$\mathbf{B}^T \Phi_0 (\boldsymbol{\alpha} - \boldsymbol{\beta}) = \mathbf{0}$$

where $\mathbf{K}'_0 := \Phi_0^T \mathbf{K}' \Phi_0$. Except for the last constraint and the usage of $\mathbf{K}'_0$, (18) is identical to (12). Similar to (13a), the

Algorithm 2: Semiparametric batch PSD map estimator

- 1: **Input:** $\{(\mathbf{x}_n, \phi(\mathbf{x}_n), \hat{\pi}_Q(\mathbf{x}_n))\}_{n=1}^N, \{\phi_m(f)\}_{m=1}^M, \epsilon$
- 2: **Parameters:** $\lambda, \mathbf{K}'(\mathbf{x}, \mathbf{x}'), \{\mathbf{B}_\nu(\mathbf{x})\}_{\nu=1}^{N_B}$
- 3: $\Phi = [\phi(\mathbf{x}_1), \dots, \phi(\mathbf{x}_N)]$
- 4: $\Phi_0 = \mathbf{I}_N \odot \Phi$
- 5: Form $\mathbf{K}'$, whose (n, n') -th block is $\mathbf{K}'(\mathbf{x}_n, \mathbf{x}_{n'})$
- 6: $\mathbf{K}'_0 = \Phi_0^T \mathbf{K}' \Phi_0$
- 7: $\mathbf{y} = [y(\mathbf{x}_1), \dots, y(\mathbf{x}_N)]^T$
- 8: Form $\mathbf{B}$, whose (n, ν) -th block is $B_\nu(\mathbf{x}_n)$
- 9: Obtain $(\hat{\boldsymbol{\alpha}}, \hat{\boldsymbol{\beta}})$ from (18) and set $\hat{\boldsymbol{\theta}} := [\hat{\boldsymbol{\theta}}_1^T, \dots, \hat{\boldsymbol{\theta}}_{N_B}^T]^T$ to be the optimal Lagrange multiplier of the last constraint
- 10: $[\hat{\mathbf{c}}_1^T, \dots, \hat{\mathbf{c}}_N^T]^T = [1/(2\lambda N)] \Phi_0 (\hat{\boldsymbol{\alpha}} - \hat{\boldsymbol{\beta}})$
- 11: **Output:** Function $\hat{\Gamma}(\mathbf{x}, f) = \sum_{m=1}^M \hat{l}_m(\mathbf{x}) \phi_m(f)$, where $[\hat{l}_1(\mathbf{x}), \dots, \hat{l}_M(\mathbf{x})]^T = \sum_{n=1}^N \mathbf{K}'(\mathbf{x}, \mathbf{x}_n) \hat{\mathbf{c}}_n + \sum_{\nu=1}^{N_B} \mathbf{B}_\nu(\mathbf{x}) \hat{\boldsymbol{\theta}}_\nu$.

primal vector of the nonparametric component can be readily obtained from the KKT conditions as

$$\hat{\mathbf{c}}' = \frac{1}{2\lambda N} \Phi_0 (\hat{\boldsymbol{\alpha}} - \hat{\boldsymbol{\beta}}). \quad (19)$$

Although $\boldsymbol{\theta}$ can also be obtained from the KKT conditions, this approach is numerically unstable. It is preferable to obtain $\boldsymbol{\theta}$ from the Lagrange multipliers of (18), which are known e.g. if an interior-point solver is employed. Specifically, noting that (17) is the dual of (18), it can be seen that $\hat{\boldsymbol{\theta}}$ equals the multipliers of the last constraint in (18). Algorithm 2 summarizes the proposed semiparametric estimator.

D. Regression with conditionally positive definite kernels

So far, the kernels were required to be positive definite. This section extends the semiparametric estimator in Sec. III-C to accommodate the wider class of *conditionally positive definite* (CPD) kernels. CPD kernels are natural for estimation problems that are invariant to translations in the data [32, p. 52], as occurs in spectrum cartography. Accommodating CPD kernels also offers a generalization of thin-plate splines (TPS), which have well-documented merits in capturing shadowing of propagation channels [8], to operate on quantized data.

Consider the following definition, which generalizes that of scalar CPD kernels [32, Sec. 2.4]. Recall that, given $\{\mathbf{x}_n\}_{n=1}^N$, $\mathbf{K}$ is a matrix whose (n, n') -th block is $\mathbf{K}(\mathbf{x}_n, \mathbf{x}_{n'})$.

Definition 1. A kernel $\mathbf{K}(\mathbf{x}_1, \mathbf{x}_2) : \mathbb{R}^d \times \mathbb{R}^d \rightarrow \mathbb{R}^M$ is CPD with respect to $\{\mathbf{B}_\nu(\mathbf{x})\}_{\nu=1}^{N_B}$ if it satisfies $\mathbf{c}^T \mathbf{K} \mathbf{c} \geq 0$ for every finite set $\{\mathbf{x}_n\}_{n=1}^N$ and all $\mathbf{c}$ such that $\mathbf{B}^T \mathbf{c} = \mathbf{0}$.

Observe that any positive definite kernel is also CPD since it satisfies $\mathbf{c}^T \mathbf{K} \mathbf{c} \geq 0$ for all $\mathbf{c}$.

To see how CPD kernels can be applied in semiparametric regression, note that (19) together with the last constraint in (18) imply that the solution to (17) satisfies $\mathbf{B}^T \hat{\mathbf{c}}' = \mathbf{0}$. Therefore, (17) can be equivalently solved by confining the vectors $\mathbf{c}'$ to lie in the null space of $\mathbf{B}^T$:

$$(\hat{\mathbf{c}}', \hat{\boldsymbol{\theta}}, \hat{\boldsymbol{\xi}}, \hat{\boldsymbol{\zeta}}) = \arg \min_{\mathbf{c}', \boldsymbol{\theta}, \boldsymbol{\xi}, \boldsymbol{\zeta}} \mathbf{1}_N^T (\boldsymbol{\xi} + \boldsymbol{\zeta}) + \lambda N \mathbf{c}'^T \mathbf{K}' \mathbf{c}'$$

s.t. $\boldsymbol{\xi} \geq \mathbf{y} - \Phi_0^T (\mathbf{K}' \mathbf{c}' + \mathbf{B} \boldsymbol{\theta}) - \epsilon \mathbf{1}_N$, $\boldsymbol{\xi} \geq \mathbf{0}_N$

$$\zeta \geq -\mathbf{y} + \Phi_0^T(\mathbf{K}'\mathbf{c}' + \mathbf{B}\boldsymbol{\theta}) - \epsilon\mathbf{1}_N, \quad \zeta \geq \mathbf{0}_N$$

$$\mathbf{B}^T\mathbf{c}' = \mathbf{0}. \quad (20)$$

The new equality constraint ensures that the objective of (20) is convex in the feasible set if $\mathbf{K}'(\mathbf{x}, \mathbf{x}')$ is CPD with respect to $\{\mathbf{B}_\nu(\mathbf{x})\}_{\nu=1}^{N_B}$. However, (20) is susceptible to numerical issues because the block matrix $\mathbf{K}'$ may not be positive semidefinite. One can circumvent this difficulty by adopting a change of variables $\mathbf{c}' := \mathbf{P}_B^\perp \tilde{\mathbf{c}}$, where $\mathbf{P}_B^\perp := \mathbf{I}_{MN} - \mathbf{B}(\mathbf{B}^T\mathbf{B})^{-1}\mathbf{B}^T \in \mathbb{R}^{MN \times MN}$ is the orthogonal projector onto the null space of $\mathbf{B}^T$. By doing so, (20) becomes

$$(\hat{\mathbf{c}}, \hat{\boldsymbol{\theta}}, \hat{\boldsymbol{\xi}}, \hat{\zeta}) = \arg \min_{\tilde{\mathbf{c}}, \boldsymbol{\theta}, \boldsymbol{\xi}, \zeta} \mathbf{1}_N^T(\boldsymbol{\xi} + \zeta) + \lambda N \tilde{\mathbf{c}}^T \mathbf{P}_B^\perp \mathbf{K}' \mathbf{P}_B^\perp \tilde{\mathbf{c}}$$

$$\text{s.t. } \boldsymbol{\xi} \geq \mathbf{y} - \Phi_0^T(\mathbf{K}' \mathbf{P}_B^\perp \tilde{\mathbf{c}} + \mathbf{B}\boldsymbol{\theta}) - \epsilon\mathbf{1}_N, \quad \boldsymbol{\xi} \geq \mathbf{0}_N$$

$$\zeta \geq -\mathbf{y} + \Phi_0^T(\mathbf{K}' \mathbf{P}_B^\perp \tilde{\mathbf{c}} + \mathbf{B}\boldsymbol{\theta}) - \epsilon\mathbf{1}_N, \quad \zeta \geq \mathbf{0}_N \quad (21)$$

where $\mathbf{P}_B^\perp \mathbf{K}' \mathbf{P}_B^\perp$ is guaranteed to be positive semidefinite.

A similar argument applies to the dual formulation in (18), where, for feasible $\boldsymbol{\alpha}$ and $\boldsymbol{\beta}$, it holds that $(\boldsymbol{\alpha} - \boldsymbol{\beta})^T \mathbf{K}'_0(\boldsymbol{\alpha} - \boldsymbol{\beta}) \geq 0$ with $\mathbf{K}'(\mathbf{x}, \mathbf{x}')$ CPD. To avoid numerical issues, observe that any feasible $\boldsymbol{\alpha}, \boldsymbol{\beta}$ satisfy $\Phi_0(\boldsymbol{\alpha} - \boldsymbol{\beta}) = \mathbf{P}_B^\perp \Phi_0(\boldsymbol{\alpha} - \boldsymbol{\beta})$, and thus (18) can be equivalently expressed as

$$(\hat{\boldsymbol{\alpha}}, \hat{\boldsymbol{\beta}}) = \arg \min_{\boldsymbol{\alpha}, \boldsymbol{\beta}} \frac{1}{4N\lambda} (\boldsymbol{\alpha} - \boldsymbol{\beta})^T \tilde{\mathbf{K}} (\boldsymbol{\alpha} - \boldsymbol{\beta})$$

$$- (\mathbf{y} - \epsilon\mathbf{1}_N)^T \boldsymbol{\alpha} + (\mathbf{y} + \epsilon\mathbf{1}_N)^T \boldsymbol{\beta} \quad (22)$$

$$\text{s. to } \mathbf{B}^T \Phi_0(\boldsymbol{\alpha} - \boldsymbol{\beta}) = \mathbf{0}$$

$$\mathbf{0}_N \leq \boldsymbol{\alpha} \leq \mathbf{1}_N, \mathbf{0}_N \leq \boldsymbol{\beta} \leq \mathbf{1}_N$$

where $\mathbf{K}'_0$ has been replaced with the positive semidefinite matrix $\tilde{\mathbf{K}} := \Phi_0^T \mathbf{P}_B^\perp \mathbf{K}' \mathbf{P}_B^\perp \Phi_0$. With this reformulation, although the optimal $\hat{\mathbf{c}}'$ can still be recovered from (19), $\hat{\boldsymbol{\theta}}$ can no longer be obtained as the Lagrange multiplier $\boldsymbol{\mu}$ of the equality constraint in (22). This is due to the change of variables, which alters (22) from being the dual of (21). As derived in Appendix A, $\hat{\boldsymbol{\theta}}$ can instead be recovered as

$$\hat{\boldsymbol{\theta}} = \boldsymbol{\mu} - (\mathbf{B}^T \mathbf{B})^{-1} \mathbf{B}^T \mathbf{K}' \hat{\mathbf{c}}'. \quad (23)$$

Broadening the scope of semiparametric regression to include CPD kernels leads also to generalizations of TPS – arguably the most popular semiparametric interpolator, which derives its name because the TPS estimate mimics the shape of a thin metal plate that minimizes the bending energy when anchored to the data points [10]. With $\mathbf{x} := [x_1, \dots, x_d]^T$, TPS adopts the parametric basis $\mathbf{B}_1(\mathbf{x}) = \mathbf{I}_M$, $\mathbf{B}_2(\mathbf{x}) = x_1 \mathbf{I}_M, \dots, \mathbf{B}_{1+d}(\mathbf{x}) = x_d \mathbf{I}_M$, and the diagonal matrix kernel

$$\mathbf{K}'(\mathbf{x}_1, \mathbf{x}_2) = r(\|\mathbf{x}_1 - \mathbf{x}_2\|_2^2) \mathbf{I}_M \quad (24)$$

where $r(z)$ denotes the radial basis function

$$r(z) := \begin{cases} z^{2s-d} \log(z) & \text{if } d \text{ is even} \\ z^{2s-d} & \text{otherwise} \end{cases} \quad (25)$$

for a positive integer s typically set to $s = 2$ [39, eq. (2.4.9)], [8]. The kernel in (24) can be shown to be CPD with respect to $\{\mathbf{B}_\nu(\mathbf{x})\}_{\nu=1}^{1+d}$ [39, p. 32]. The norm in the RKHS

$\mathcal{H}'$ induced by (24), which can be evaluated as in Sec. III-A, admits the equivalent form

$$\|\mathbf{l}'\|_{\mathcal{H}'}^2 = \sum_{m=1}^M \int_{\mathbb{R}^d} \|\nabla^2 l_m(\mathbf{z})\|_F^2 d\mathbf{z}$$

where $\|\cdot\|_F$ denotes Frobenius norm, and ∇^2 the Hessian. Therefore, the RKHS norm of $\mathcal{H}'$ captures the conventional notion of smoothness embedded in the magnitude of the second-order derivatives. Among other reasons, TPS are popular because they do not require parameter tuning, unlike e.g., Gaussian kernels, which need adjustment of their variance parameter. The novelty here is the generalization of TPS to vector-valued function estimation from quantized observations. Unlike [8], which relies on un-quantized periodograms, the proposed scheme is based on quantized power measurements.

E. Multiple measurements per sensor

For simplicity, it was assumed so far that each sensor collects and reports a single measurement to the fusion center. However, $P > 1$ measurements can be obtained per sensor by changing the filter impulse response between measurements, or, by appropriately modifying the compression matrix in their A2ICs; cf. Remark 3. A naive approach would be to regard the P measurements per sensor as measurements from P different sensors at the same location. However, this increases the problem size by a factor of P , and one has to deal with a rank deficient kernel matrix $\mathbf{K}$ of dimension MNP , which renders the solutions of (7), (10) and (16) non-unique.

A more efficient means of accommodating P measurements per sensor in the un-quantized scenario is to reformulate (5) as

$$\hat{\mathbf{l}} = \arg \min_{\mathbf{l} \in \mathcal{H}} \frac{1}{NP} \sum_{n=1}^N \sum_{p=1}^P \mathcal{L}(\hat{\pi}_p(\mathbf{x}_n) - \phi_p^T(\mathbf{x}_n) \mathbf{l}(\mathbf{x}_n)) + \lambda \|\mathbf{l}\|_{\mathcal{H}}^2 \quad (26)$$

where $\hat{\pi}_p(\mathbf{x}_n)$ and $\phi_p(\mathbf{x}_n)$ correspond to the p -th measurement reported by the sensor at location $\mathbf{x}_n$. Collect all observations in $\hat{\boldsymbol{\pi}} := [\hat{\pi}_1(\mathbf{x}_1), \hat{\pi}_2(\mathbf{x}_1), \dots, \hat{\pi}_P(\mathbf{x}_N)]^T$, let $\boldsymbol{\Phi} := [\phi_1(\mathbf{x}_1), \phi_2(\mathbf{x}_1), \dots, \phi_P(\mathbf{x}_N)]$, and let $\bar{\boldsymbol{\Phi}}_0 := (\mathbf{I}_N \otimes \mathbf{1}_P^T) \odot \boldsymbol{\Phi} \in \mathbb{R}^{MN \times NP}$. As before, the representer theorem implies that the minimizer of (26) is given by $\hat{\mathbf{l}}(\mathbf{x}) = \sum_{n=1}^N \mathbf{K}(\mathbf{x}, \mathbf{x}_n) \hat{\mathbf{c}}_n := \mathbf{K}(\mathbf{x}) \hat{\mathbf{c}}$. If $\mathcal{L}$ is the square loss $\mathcal{L}_2$ (see Sec. III-A), then

$$\hat{\mathbf{c}} = (\bar{\boldsymbol{\Phi}}_0 \bar{\boldsymbol{\Phi}}_0^T \mathbf{K} + \lambda N \mathbf{I}_{MN})^{-1} \bar{\boldsymbol{\Phi}}_0 \hat{\boldsymbol{\pi}} \quad (27)$$

which generalizes (8) to $P \geq 1$. Note however that $\bar{\mathbf{c}}$ has dimension MN , whereas the aforementioned naive approach would result in MNP . Likewise, $\mathbf{K}$ is of size $MN \times MN$.

If $\hat{\pi}_p(\mathbf{x}_n)$ is replaced with $y_p(\mathbf{x}_n)$ in (26), the resulting expression extends (9) to multiple measurements per sensor. If $\mathcal{L}$ is the ϵ -insensitive cost $\mathcal{L}_{1\epsilon}$, such an expression is minimized for $\hat{\mathbf{l}}(\mathbf{x}) = \mathbf{K}(\mathbf{x}) \hat{\mathbf{c}}$ with

$$(\hat{\mathbf{c}}, \hat{\boldsymbol{\xi}}, \hat{\zeta}) = \arg \min_{\tilde{\mathbf{c}}, \boldsymbol{\xi}, \zeta} \mathbf{1}_{NP}^T(\bar{\boldsymbol{\xi}} + \bar{\zeta}) + \lambda NP \tilde{\mathbf{c}}^T \mathbf{K} \tilde{\mathbf{c}}$$

$$\text{s. to } \bar{\boldsymbol{\xi}} \geq \bar{\mathbf{y}} - \bar{\boldsymbol{\Phi}}_0^T \mathbf{K} \tilde{\mathbf{c}} - \epsilon \mathbf{1}_{NP}, \quad \bar{\boldsymbol{\xi}} \geq \mathbf{0}_{NP}$$

$$\bar{\zeta} \geq -\bar{\mathbf{y}} + \bar{\boldsymbol{\Phi}}_0^T \mathbf{K} \tilde{\mathbf{c}} - \epsilon \mathbf{1}_{NP}, \quad \bar{\zeta} \geq \mathbf{0}_{NP} \quad (28)$$

where $\bar{\mathbf{y}} := [y_1(\mathbf{x}_1), y_2(\mathbf{x}_1), \dots, y_P(\mathbf{x}_N)]^T$. Note that (28) reduces to (11) if $P = 1$. The dual formulation is the same as (12), except that α, β, Φ_0 , and N are replaced with $\bar{\alpha}, \bar{\beta}, \bar{\Phi}_0$, and NP , respectively. The primal solution can be recovered as $\hat{\mathbf{c}} = (2\lambda NP)^{-1} \bar{\Phi}_0(\bar{\alpha} - \bar{\beta})$; cf. (13a).

The multi-measurement version of the semiparametric estimator in (15) is

$$\hat{\mathbf{l}} = \arg \min_{\mathbf{l} \in \mathcal{H}} \frac{1}{NP} \sum_{n=1}^N \sum_{p=1}^P \mathcal{L}(y_p(\mathbf{x}_n) - \phi_p^T(\mathbf{x}_n) \mathbf{l}(\mathbf{x}_n)) + \lambda \|\mathbf{l}'\|_{\mathcal{H}}^2 \quad (29)$$

and the counterpart to (17) is obtained by replacing $\xi, \zeta, \mathbf{y}, \phi(\mathbf{x}_n)$ and N , with $\bar{\xi}, \bar{\zeta}, \bar{\mathbf{y}}, \phi_p(\mathbf{x}_n)$ and NP , respectively. Likewise, the dual formulation is obtained by substituting α, β, Φ_0, N , and $\bar{\mathbf{K}}$ in (22) with $\bar{\alpha}, \bar{\beta}, \bar{\Phi}_0, NP$, and

$$\bar{\mathbf{K}} := \bar{\Phi}_0^T \mathbf{P}_B^\perp \mathbf{K}' \mathbf{P}_B^\perp \bar{\Phi}_0 \quad (30)$$

respectively. The primal variables are recovered again as $\hat{\mathbf{c}} = (2\lambda NP)^{-1} \bar{\Phi}_0(\bar{\alpha} - \bar{\beta})$ and $\hat{\boldsymbol{\theta}} = \bar{\boldsymbol{\mu}} - (\mathbf{B}^T \mathbf{B})^{-1} \mathbf{B}^T \mathbf{K}' \hat{\mathbf{c}}$, where $\bar{\boldsymbol{\mu}}$ is the Lagrange multiplier vector associated with the equality constraints in the dual problem; cf. (23).

Remark 5 (Non-uniform quantization). Unless $\hat{\pi}(\mathbf{x}_n)$ is uniformly distributed, non-uniform quantization may be preferable over the uniform quantization adopted so far. With R quantization regions specified by the boundaries $0 = \tau_0 < \tau_1 < \dots < \tau_R$, the quantized measurements are $\pi_Q(\mathbf{x}_n) := Q(\hat{\pi}(\mathbf{x}_n)) = i$ for $\hat{\pi}(\mathbf{x}_n) \in [\tau_i, \tau_{i+1})$. The general formulations (26) and (29) can accommodate non-uniformly quantized observations by replacing ϵ in all relevant optimization problems with $\epsilon(\mathbf{x}_n) := (\tau_{\pi_Q(\mathbf{x}_n)+1} - \tau_{\pi_Q(\mathbf{x}_n)})/2$; and likewise modifying the centroid expression from $\mathbf{y}(\mathbf{x}_n) := [2\pi_Q(\mathbf{x}_n) + 1]\epsilon$ to $\mathbf{y}(\mathbf{x}_n) := (\tau_{\pi_Q(\mathbf{x}_n)+1} + \tau_{\pi_Q(\mathbf{x}_n)})/2$.

Remark 6 (Enforcing nonnegativity). Since all M entries of vector $\mathbf{l}(\mathbf{x})$ represent power, they are inherently nonnegative. To exploit this information, at least partially, one can enforce non-negativity of $\mathbf{l}(\mathbf{x})$ at all sensor locations by introducing the constraint $\mathbf{K}(\mathbf{x}_n) \mathbf{c} \geq \mathbf{0}_M$ for $n = 1, \dots, N$ in (11). Another approach is to include M “virtual measurements” for every sensor location $\mathbf{x}_n$ to promote estimates $\mathbf{l}(\mathbf{x})$ satisfying $0 \leq l_m(\mathbf{x}_n) < \tau_R$ for all m . In this way, the estimation algorithm does not use the set of “real” measurements $\{(y_p(\mathbf{x}_n), \phi_p(\mathbf{x}_n), \epsilon_p(\mathbf{x}_n))\}_{p=1}^P$ for every sensor $n = 1, \dots, N$, where $\epsilon_p(\mathbf{x}_n)$ is the quantization interval of the p -th measurement obtained by the n -th sensor. Instead, it uses its extended version $\{(y_p(\mathbf{x}_n), \phi_p(\mathbf{x}_n), \epsilon_p(\mathbf{x}_n))\}_{p=1}^{P+M}$, where $\phi_{P+m}(\mathbf{x}_n) = \mathbf{e}_{M,m}$, $y_{P+m}(\mathbf{x}_n) := (\tau_0 + \tau_R)/2$, $\epsilon_{P+m}(\mathbf{x}_n) := (\tau_R - \tau_0)/2$ for $m = 1, \dots, M$. The measurements $p = P+1, \dots, P+M$ are termed “virtual” since they are appended to the “real” measurements by the fusion center, but are not acquired by the sensors. Note that this approach does not constrain $\mathbf{l}$ to be entry-wise non-negative at the sensor locations, it just promotes estimates satisfying this condition.

Remark 7 (Computational complexity). With un-quantized data, the estimate (27) requires $\mathcal{O}(M^3 N^3 + PM^2 N)$ operations. With nonparametric estimation from quantized data,

solving the dual of (28) through interior point methods takes $\mathcal{O}((NP)^{3.5})$ iterations. A similar level of complexity is incurred by its semiparametric counterpart. Albeit polynomial, this complexity may be prohibitive in real-time applications with limited computational capabilities if the number of measurements NP is large. For such scenarios, an online algorithm with linear complexity is proposed in Sec. IV, which is guaranteed to converge to the nonparametric estimate (10). An online algorithm for semiparametric estimation can be found in [30].

Remark 8. In real applications, low SNR conditions, transmission beamforming, and the hidden terminal problem may limit the quality of PSD map estimates relying on short observation windows. Thus, highly accurate estimates may require longer observation intervals to average out the undesirable effects of noise and small-scale fading. This constitutes a tradeoff between estimation accuracy and temporal resolution of PSD maps that is inherent to the spectrum cartography problem. This tradeoff may not pose a difficulty in applications such as in TV networks, where transmitters remain active or inactive for long time intervals, typically months or years, and therefore a high temporal resolution is unnecessary. In other scenarios, it may suffice to know whether the primary users are active or inactive, in which case a high estimation accuracy is not needed, and therefore short observation windows may be enough. However, in those scenarios where both high accuracy and fine temporal resolution are required, one would need to deploy more sensors. The present paper alleviates the negative impact of the aforementioned tradeoff by reducing the required communication bandwidth. This brings a threefold benefit: (i) in a fixed time interval, each sensor may report more measurements to the fusion center, thus improving averaging and therefore the accuracy of the estimates; (ii), a larger number of sensors can be deployed; and (iii), the latency of the communication between sensors and fusion center is reduced.

IV. ONLINE ALGORITHM

The algorithms proposed so far operate in batch mode, thus requiring all observations before they can commence. Moreover, their computational complexity increases faster than linearly in NP , which may be prohibitive if the number of measurements is large relative to the available computational resources. These considerations motivate the development of online algorithms, which can both approximate the solution of the batch problem with complexity $\mathcal{O}(NP)$ and update $\hat{\mathbf{l}}(\mathbf{x})$ as new measurements arrive at the fusion center. Online algorithms are further motivated by their ability to track a time-varying $\mathbf{l}$.

Although online algorithms can be easily constructed by iteratively applying batch algorithms over sliding windows [33], online strategies with instantaneous updates are preferred [14]. An elegant approach for kernel-based learning relying on stochastic gradient descent (SGD) in the RKHS is developed in [24] for scalar kernel machines. Its counterpart for vector-valued functions is described in [4], but it is not directly

applicable to the present setup since it requires differentiable objectives and vectors $\phi(\mathbf{x}_\nu)$ that do not depend on ν . This section extends the algorithm in [4] to accommodate the present scenario.

For a selected $\mathcal{L}$, consider the instantaneous regularized cost, defined for generic $\mathbf{l} \in \mathcal{H}$, $\mathbf{x}_\nu$, $\phi(\mathbf{x}_\nu)$, and $y(\mathbf{x}_\nu)$ as

$$\mathcal{C}(\mathbf{l}, \phi(\mathbf{x}_\nu), \mathbf{x}_\nu, y(\mathbf{x}_\nu)) := \mathcal{L}(y(\mathbf{x}_\nu) - \phi^T(\mathbf{x}_\nu)\mathbf{l}(\mathbf{x}_\nu)) + \lambda \|\mathbf{l}\|_{\mathcal{H}}^2. \quad (31)$$

Note that (5) indeed minimizes the sample average of $\mathcal{C}$. Suppose that at time index $t = 1, 2, \dots$ the fusion center processes one measurement from sensor $n_t \in \{1, \dots, N\}$. If the fusion center uses multiple observations, say P , from the n' -th sensor, then $n_{t_1} = n_{t_2} = \dots = n_{t_P} = n'$, where $\{t_1, \dots, t_P\}$ depends on the fusion center schedule.

Upon processing the t -th measurement, the SGD update is

$$\mathbf{l}^{(t+1)}(\mathbf{x}) = \mathbf{l}^{(t)}(\mathbf{x}) - \mu_t \partial_{\mathbf{l}} \mathcal{C}(\mathbf{l}^{(t)}, \phi(\mathbf{x}_{n_t}), \mathbf{x}_{n_t}, y(\mathbf{x}_{n_t}))(\mathbf{x}) \quad (32)$$

where $\mathbf{l}^{(t)}$ is the estimate at time t before $y(\mathbf{x}_{n_t})$ has been received; $\mu_t > 0$ is the learning rate; and $\partial_{\mathbf{l}}$ denotes subgradient with respect to $\mathbf{l}$. In general, μ_t can be replaced with a matrix $\mathbf{M}_t$ to increase flexibility. For $\mathcal{C}$ as in (31), it can be shown using [4, eq. (2)] that

$$\begin{aligned} \partial_{\mathbf{l}} \mathcal{C}(\mathbf{l}, \phi(\mathbf{x}_\nu), \mathbf{x}_\nu, y(\mathbf{x}_\nu))(\mathbf{x}) \\ = -\mathcal{L}'(y(\mathbf{x}_\nu) - \phi^T(\mathbf{x}_\nu)\mathbf{l}(\mathbf{x}_\nu))\mathbf{K}(\mathbf{x}, \mathbf{x}_\nu)\phi(\mathbf{x}_\nu) + 2\lambda\mathbf{l}(\mathbf{x}) \end{aligned} \quad (33)$$

and $\mathcal{L}'$ is a subgradient of $\mathcal{L}$, which for $\mathcal{L}_{1\epsilon}(e_\nu) := \max(0, |e_\nu| - \epsilon)$ is e.g.:

$$\mathcal{L}'_{1\epsilon}(e_\nu) = \frac{\text{sgn}(e_\nu - \epsilon) + \text{sgn}(e_\nu + \epsilon)}{2}. \quad (34)$$

Substituting (33) into (32) yields

$$\begin{aligned} \mathbf{l}^{(t+1)}(\mathbf{x}) &= (1 - 2\mu_t\lambda)\mathbf{l}^{(t)}(\mathbf{x}) \\ &+ \mu_t\mathcal{L}'(y(\mathbf{x}_{n_t}) - \phi^T(\mathbf{x}_{n_t})\mathbf{l}^{(t)}(\mathbf{x}_{n_t}))\mathbf{K}(\mathbf{x}, \mathbf{x}_{n_t})\phi(\mathbf{x}_{n_t}). \end{aligned} \quad (35)$$

Upon setting $\mathbf{l}^{(1)} = \mathbf{0}$, it follows that

$$\mathbf{l}^{(t)}(\mathbf{x}) = \sum_{i=1}^{t-1} \mathbf{K}(\mathbf{x}, \mathbf{x}_{n_i})\mathbf{c}_i^{(t)} \quad (36)$$

for some $\mathbf{c}_i^{(t)}$, $i = 1, \dots, t-1$. Interestingly, although the representer theorem [26, Thm. 5] has not been invoked, the estimates $\mathbf{l}^{(t)}(\mathbf{x})$ here have the form of those in Sec. III.

If measurements come from sensors at different locations, the functions $\mathbf{K}(\mathbf{x}, \mathbf{x}_{n_i})$, $i = 1, \dots, t$, are linearly independent for properly selected kernels, and substituting (36) into (35) results in the following update rule:

$$\mathbf{c}_i^{(t+1)} = \begin{cases} (1 - 2\mu_t\lambda)\mathbf{c}_i^{(t)} & \text{if } i = 1, \dots, t-1 \\ \mu_t\mathcal{L}'[y(\mathbf{x}_{n_t}) - \phi^T(\mathbf{x}_{n_t})\mathbf{l}^{(t)}(\mathbf{x}_{n_t})]\phi(\mathbf{x}_{n_t}) & \text{if } i = t. \end{cases}$$

This equation reveals that the number of coefficients maintained increases linearly in t . This is the so-called *curse of kernelization* [41]. However, if $\mu_t\lambda \in (0, 1)$, then the

amplitudes of the entries in $\mathbf{c}_i^{(t)}$ are shrunk by a factor $|1 - 2\mu_t\lambda| < 1$. This justifies truncating (36) as

$$\mathbf{l}^{(t)}(\mathbf{x}) = \sum_{i=\max(1, t-I)}^{t-1} \mathbf{K}(\mathbf{x}, \mathbf{x}_i)\mathbf{c}_i^{(t)} \quad (37)$$

for some $I > 1$. On the other hand, if the fusion center processes multiple observations per sensor, $\{\mathbf{K}(\mathbf{x}, \mathbf{x}_{n_i})\}_{i=1}^t$ are no longer linearly independent. In such a case, up to N kernels $\{\mathbf{K}(\mathbf{x}, \mathbf{x}_n)\}_{n=1}^N$ are linearly independent, which yields

$$\mathbf{l}^{(t)}(\mathbf{x}) = \sum_{n=1}^N \mathbf{K}(\mathbf{x}, \mathbf{x}_n)\mathbf{c}_n^{(t)}. \quad (38)$$

After receiving the observation $(\mathbf{x}_{n_t}, y(\mathbf{x}_{n_t}), \phi(\mathbf{x}_{n_t}))$ at time t , it follows from (38) that one must obtain $\{\mathbf{c}_n^{(t+1)}\}_{n=1}^N$ as

$$\mathbf{c}_n^{(t+1)} = \begin{cases} (1 - 2\mu_t\lambda)\mathbf{c}_n^{(t)} & \text{if } n \neq n_t \\ (1 - 2\mu_t\lambda)\mathbf{c}_n^{(t)} + \mu_t\mathcal{L}'[y(\mathbf{x}_{n_t}) - \phi^T(\mathbf{x}_{n_t})\mathbf{l}^{(t)}(\mathbf{x}_{n_t})]\phi(\mathbf{x}_{n_t}) & \text{if } n = n_t. \end{cases}$$

Convergence of these recursions is characterized by the next result, which adapts [4, Thm. 1] to the proposed setup.

Theorem 1. *If $\lambda_{\max}(\mathbf{K}(\mathbf{x}, \mathbf{x})) < \bar{\lambda}^2 < \infty$ for all $\mathbf{x}$, $\|\phi(\mathbf{x}_{n_t})\|_2 \leq \bar{\varphi}$ for all t , and $\mu_t := \mu t^{-1/2}$ with $\mu\lambda < 1$, then the iterates in (35) with $\mathcal{L} = \mathcal{L}_{1\epsilon}$ satisfy*

$$\begin{aligned} \frac{1}{T} \sum_{t=1}^T \mathcal{C}(\mathbf{l}^{(t)}, \phi(\mathbf{x}_{n_t}), \mathbf{x}_{n_t}, y(\mathbf{x}_{n_t})) \\ \leq \inf_{\mathbf{l} \in \mathcal{H}} \left[\frac{1}{T} \sum_{t=1}^T \mathcal{C}(\mathbf{l}, \phi(\mathbf{x}_{n_t}), \mathbf{x}_{n_t}, y(\mathbf{x}_{n_t})) \right] + \frac{c_1}{\sqrt{T}} + \frac{c_2}{T} \end{aligned} \quad (39)$$

where $c_2 := \bar{\lambda}^2 \bar{\varphi}^2 / (8\lambda^2 \mu)$ and $c_1 := 4(\bar{\lambda}^2 \bar{\varphi}^2 \mu + c_2)$.

Proof: See Appendix C. ■

In words, Theorem 1 establishes that the averaged instantaneous error from the online algorithm converges to the regularized empirical error of the batch solution.

V. NUMERICAL TESTS

In this section, the proposed algorithms are validated through numerical experiments. Following [17], a correlated shadow fading model was adopted, where, for $m = 1, \dots, M-1$,

$$10 \log l_m(\mathbf{x}) = 10 \log A_m - \gamma \log(\delta + \|\mathbf{x} - \mathbf{x}_m\|) + s_m(\mathbf{x}). \quad (40)$$

Here, $\delta > 0$ is a small constant ensuring that the argument of the logarithm does not vanish, $\gamma = 3$ is the pathloss exponent, and the parameters A_m and $\mathbf{x}_m$ denote the power and location of the m -th transmitter, respectively. The random shadowing component $s_m(\mathbf{x})$ is generated as a zero-mean Gaussian random variable with $\mathbb{E}\{s_m(\mathbf{x})s_m(\mathbf{x}')\} = \sigma_s^2 \rho^{-\|\mathbf{x} - \mathbf{x}'\|}$, where $\sigma_s^2 = 2$ and $\rho = 0.8$. The noise power was set to $l_M(\mathbf{x}) = 0.75$.

The N sensors, deployed uniformly at random over the region of interest, report P quantized measurements $\{\hat{\pi}_{Q,p}(\mathbf{x}_n) = Q(|\pi_p(\mathbf{x}_n) + \eta_p(\mathbf{x}_n)|) = Q(|\phi_p^T(\mathbf{x}_n)\mathbf{l}(\mathbf{x}_n) +$

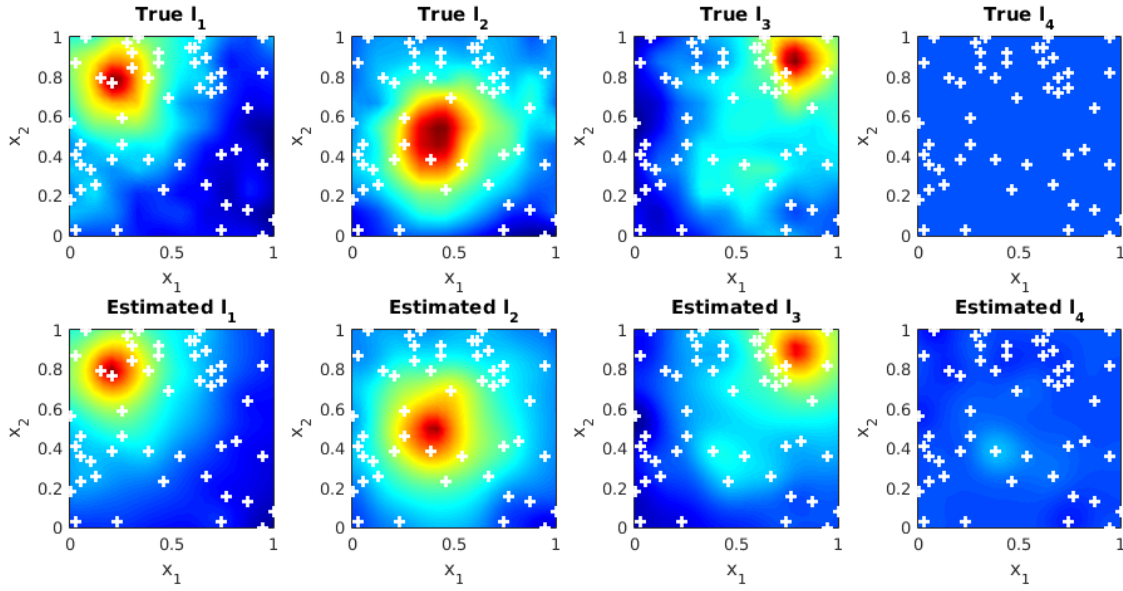


Fig. 2: True and estimated functions $l_m(\mathbf{x})$ with $N = 50$, $P = 8$, 5-bit quantization with CPQ, $\lambda = 10^{-9}$, and nonnegativity enforced through virtual sensors. Each column corresponds to the power of a channel. White crosses denote sensor locations. The PSD at any location can be reconstructed by substituting the values of these functions into (1).

$\eta_p(\mathbf{x}_n)|\}_{p=1}^P$ to the fusion center, where $\eta_p(\mathbf{x}_n) \sim \mathcal{N}(0, \sigma_\eta^2)$ simulates noise due to finite sample estimation effects. The entries of $\phi_p(\mathbf{x}_n)$ were generated uniformly over the interval $[0, 1]$ for all $\mathbf{x}_n$ and p .

Two quantization schemes were implemented. Under uniform quantization (UQ), the range of $\pi_p(\mathbf{x})$ was first determined using Monte Carlo runs and the quantization region boundaries $\tau_0 < \tau_1 < \dots < \tau_R$ were set such that $\tau_{i+1} - \tau_i = 2\epsilon \forall i$, where ϵ was such that the probability of clipping $\Pr\{\pi_p(\mathbf{x}) > \tau_R\}$ was approximately 10^{-3} and $R := 2^b$, with b the number of bits per measurement. Under constant probability quantization (CPQ), these boundaries were chosen such that $\Pr\{\tau_i \leq \pi_p(\mathbf{x}) < \tau_{i+1}\}$ was approximately constant for all i .

The true PSD map in the region $[0, 1] \times [0, 1] \subset \mathbb{R}^2$ ($d = 2$) was created with $M - 1 = 3$ transmitters, $A_1 = 0.9$, $A_2 = 0.8$, $A_3 = 0.7$, $\chi_1 = (0.2, 0.8)$, $\chi_2 = (0.4, 0.5)$, and $\chi_3 = (0.8, 0.9)$. Using CPQ and enforcing nonnegativity through M virtual measurements per sensor, the batch semiparametric estimate from (29) was computed. The nonparametric part adopts a diagonal Gaussian kernel matrix $\mathbf{K}(\mathbf{x}_n, \mathbf{x}_{n'})$ with $k_m(\mathbf{x}_n, \mathbf{x}_{n'}) = \exp(-\|\mathbf{x}_n - \mathbf{x}_{n'}\|^2 / \sigma_m^2)$ on its m -th diagonal entry (cf. Remark 4), where $\sigma_m^2 = 0.12$ for $m = 1, 2, 3, 4$. The parametric part is spanned by a basis with $N_B = 1$ and $\mathbf{B}_1(\mathbf{x})$ a diagonal matrix whose (m, m) -th entry is given by $1/(\delta + \|\mathbf{x} - \chi_m\|^\gamma)$ if $m = 1, \dots, M - 1$; and 0 if $m = M$. The variance σ_η^2 was set such that about 15% of the measurements contain errors. Fig. 2 shows the true and estimated maps for a particular realization of sensor locations $\{\mathbf{x}_n\}$ (represented by crosses), measurement vectors $\{\phi_p(\mathbf{x}_n)\}$, and measurement noise $\eta_p(\mathbf{x}_n)$. Each sensor reports $P = 8$ measurements quantized to $b = 5$ bits. Although every sensor transmits only 5 bytes, it is observed that the reconstructed PSD maps match

well with the true ones for all transmitters.

To quantify the estimation performance, the normalized mean-square error (NMSE), defined as

$$\text{NMSE} := \frac{\mathbb{E} \left\{ \|\mathbf{l}(\mathbf{x}) - \hat{\mathbf{l}}(\mathbf{x})\|_2^2 \right\}}{\mathbb{E} \left\{ \|\mathbf{l}(\mathbf{x})\|_2^2 \right\}} \quad (41)$$

was employed, where the expectation is taken with respect to the uniformly distributed $\mathbf{x}$, measurement vectors $\{\phi_p(\mathbf{x}_n)\}$, and measurement noise $\eta_p(\mathbf{x}_n)$. To focus on quantization effects, the region of interest was the one-dimensional ($d = 1$) interval $[0, 1] \subset \mathbb{R}^1$, and $M - 1 = 4$ transmitters with parameters $\chi_1 = 0.1$, $\chi_2 = 0.2$, $\chi_3 = 0.4$, $\chi_4 = 0.8$, $A_1 = 0.8$, $A_2 = 0.9$, $A_3 = 0.8$ and $A_4 = 0.7$, were considered.

To illustrate the effects of quantization, Fig. 3 depicts the NMSE as a function of N for different values of b using both nonparametric and semiparametric estimators. To capture only quantization effects, no measurement errors were introduced ($\sigma_\eta^2 = 0$), uniform quantization was used, and σ_s^2 was set to zero. The nonparametric approach used a diagonal Gaussian kernel (GK) as in Fig. 2. It can be seen that the proposed estimators are consistent in N . Although, for this particular case, TPS mostly outperforms GK-based regression, this need not hold in different scenarios since which kernel leads to the best performance depends on the propagation environment as well as on the field characteristics.

Fig. 4 depicts the NMSE for $N = 40$ sensors with a varying number of measurements P per sensor in different settings. As expected, the estimates are seen to be inconsistent as P grows since the number of sensors is fixed and there is no way to accurately estimate the PSD map far from their vicinity. It is also seen that TPS regression benefits more from incorporating non-negativity constraints than the GK-based schemes. Finally, it is observed that CPQ outperforms

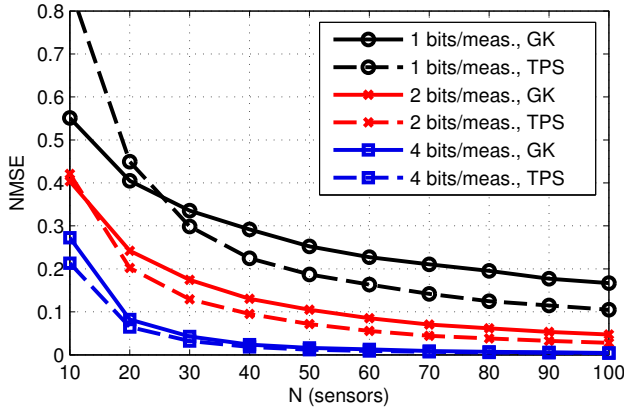


Fig. 3: NMSE versus N for variable UQ resolution with $M = 5$, $P = 5$, $\sigma_\eta^2 = 0$, $\lambda = 10^{-6}$, and nonnegativity not enforced.

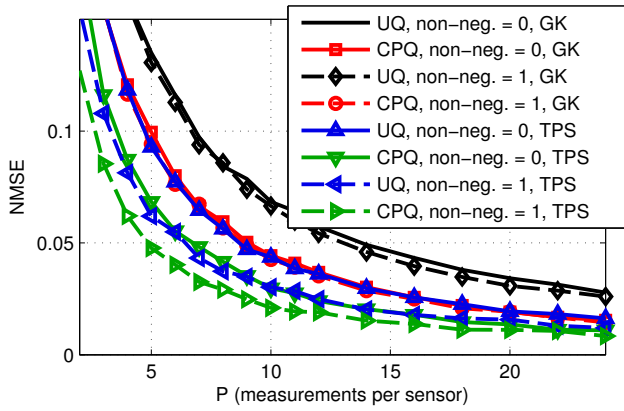


Fig. 4: Performance of different quantizers with and without nonnegativity constraints ($M = 5$, $N = 40$, $b = 2$ bits/measurement, $\sigma_\eta^2 = 0$, and $\lambda = 10^{-6}$).

UQ. However, as demonstrated by Fig. 5, UQ leads to a better performance than CPQ for this simulation scenario if the noise variance σ_η^2 is sufficiently large. Fig. 5 further shows that the effect of measurement noise is more pronounced for larger b . This is intuitive since a given measurement noise variance leads to more measurement error events. Thus, $\pi_p(x_n)$ must be estimated more accurately under finer quantization.

Fig. 6 depicts the performance of the online algorithm using the representation in (38). As a benchmark, the offline (batch) algorithm was run per time slot with all the data received up to that time slot. The top panel in Fig. 6 shows the regularized empirical risks (evaluated per time slot using the entire set of observations) for different learning rates μ_t . Common to gradient methods with constant step size, a larger μ_t speeds up convergence, but also increases the residual error. The bottom panel depicts the NMSE evolution over time. Using NMSE as figure of merit favors greater learning rates over smaller ones.

VI. CONCLUSIONS

This paper introduced a family of methods for nonparametric and semiparametric estimation of PSD maps using a set of linearly compressed and possibly quantized power measurements collected by distributed sensors. To capture

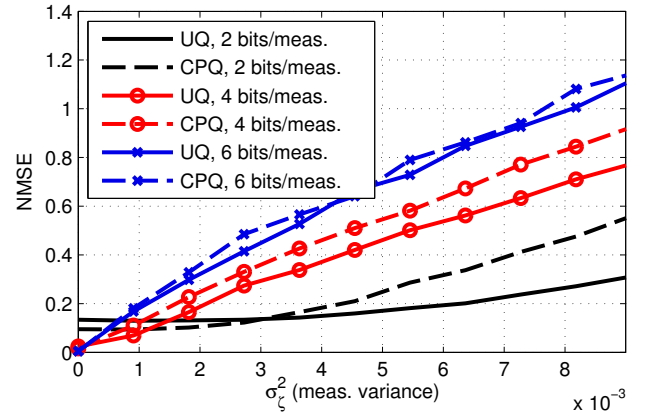


Fig. 5: Measurement noise effects ($M = 5$, $N = 40$, $P = 5$, $\lambda = 10^{-6}$, GK, and nonnegativity not enforced).

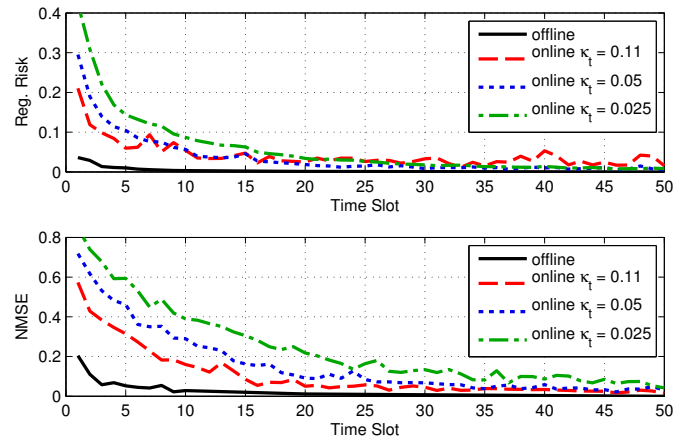


Fig. 6: Performance of the online algorithm ($M = 5$, $N = 30$, $b = 4$ bits/measurement, $\sigma_\eta^2 = 0$, $\lambda = 10^{-6}$, GK, and nonnegativity not enforced).

different degrees of prior information without sacrificing flexibility, nonparametric and semiparametric estimators have been proposed by leveraging the framework of kernel-based learning. Existing semiparametric regression techniques have been generalized to vector-valued function estimation and shown to subsume thin-plate spline regression as a special case. Extensions to multiple measurements per sensor, non-uniform quantization and non-negativity constraints have also been introduced. Both batch and online approaches were developed, thereby offering a performance complexity trade-off.

Future work will be devoted to kernel selection [16], quantizer design, and alternative types of spectral cartography formats including construction of channel-gain maps.

APPENDIX A DERIVATION OF (23)

This appendix describes how to obtain $\hat{c}$ and $\hat{\theta}$ in (21) from $(\hat{\alpha}, \hat{\beta})$ in (22). Because of the change of variable, (22) is not the dual of (21); as a result $\hat{c}$ and $\hat{\theta}$ are not, in general, the Lagrange multipliers of (22). As shown next, they can be

recovered after relating (22) to the dual of (21). The latter is:

$$\begin{aligned} (\tilde{\mathbf{c}}, \tilde{\boldsymbol{\alpha}}, \tilde{\boldsymbol{\beta}}) &= \arg \min_{\tilde{\mathbf{c}}, \tilde{\boldsymbol{\alpha}}, \tilde{\boldsymbol{\beta}}} \lambda N \tilde{\mathbf{c}}^T \mathbf{P}_B^\perp \mathbf{K}' \mathbf{P}_B^\perp \tilde{\mathbf{c}} \\ &\quad - (\mathbf{y} - \epsilon \mathbf{1}_N)^T \boldsymbol{\alpha} + (\mathbf{y} + \epsilon \mathbf{1}_N)^T \boldsymbol{\beta} \\ \text{s. to } 2\lambda N \mathbf{P}_B^\perp \mathbf{K}' \mathbf{P}_B^\perp \tilde{\mathbf{c}} - \mathbf{P}_B^\perp \mathbf{K}' \boldsymbol{\Phi}_0 (\boldsymbol{\alpha} - \boldsymbol{\beta}) &= \mathbf{0}_{MN} \\ \mathbf{B}^T \boldsymbol{\Phi}_0 (\boldsymbol{\alpha} - \boldsymbol{\beta}) &= \mathbf{0}_{MNB} \\ \boldsymbol{\alpha} - \mathbf{1}_N &\leq \mathbf{0}_N, -\boldsymbol{\alpha} \leq \mathbf{0}_N, \boldsymbol{\beta} - \mathbf{1}_N \leq \mathbf{0}_N, -\boldsymbol{\beta} \leq \mathbf{0}_N. \end{aligned} \quad (42)$$

The presence of $\tilde{\mathbf{c}}$ in both (21) and (42) is owing to the fact that $\mathbf{P}_B^\perp \mathbf{K}' \mathbf{P}_B^\perp$ is not invertible. Taking into account that (21) is indeed the dual of (42), it can be shown that $\tilde{\boldsymbol{\theta}}$ is the Lagrange multiplier $\tilde{\boldsymbol{\mu}}_2$ associated with the second equality constraint of (42), whereas $\tilde{\mathbf{c}} = \hat{\mathbf{c}}$. Thus, it remains only to express $\tilde{\boldsymbol{\mu}}_2$ and $\tilde{\mathbf{c}}$ in terms of the solution to (22).

From the second equality constraint in (42), it holds for $(\boldsymbol{\alpha}, \boldsymbol{\beta})$ feasible that $\boldsymbol{\Phi}_0 (\boldsymbol{\alpha} - \boldsymbol{\beta}) = \mathbf{P}_B^\perp \boldsymbol{\Phi}_0 (\boldsymbol{\alpha} - \boldsymbol{\beta})$. Then, the first equality constraint in (42) can be replaced with

$$2\lambda N \mathbf{P}_B^\perp \mathbf{K}' \mathbf{P}_B^\perp \tilde{\mathbf{c}} - \mathbf{P}_B^\perp \mathbf{K}' \mathbf{P}_B^\perp \boldsymbol{\Phi}_0 (\boldsymbol{\alpha} - \boldsymbol{\beta}) = \mathbf{0}_{MN}. \quad (43)$$

Solving (43) for $\tilde{\mathbf{c}}$ produces

$$\tilde{\mathbf{c}} = \frac{1}{2\lambda N} \boldsymbol{\Phi}_0 (\boldsymbol{\alpha} - \boldsymbol{\beta}) + \nu (\mathbf{P}_B^\perp \mathbf{K}' \mathbf{P}_B^\perp) \quad (44)$$

where $\nu(\mathbf{A})$ denotes any vector in the null space of $\mathbf{A}$. Substituting (44) into (42), one recovers (22).

Clearly, problems (42) and (22) are equivalent in the sense that if $(\hat{\boldsymbol{\alpha}}, \hat{\boldsymbol{\beta}})$ solves the latter, then

$$\tilde{\mathbf{c}} = \frac{1}{2\lambda N} \boldsymbol{\Phi}_0 (\hat{\boldsymbol{\alpha}} - \hat{\boldsymbol{\beta}}), \quad \tilde{\boldsymbol{\alpha}} = \hat{\boldsymbol{\alpha}}, \quad \tilde{\boldsymbol{\beta}} = \hat{\boldsymbol{\beta}} \quad (45)$$

solve the former. However, their Lagrange multipliers differ due to the transformation introduced. A possible means of establishing their relation is to compare the KKT conditions of both problems. In particular, let $\hat{\boldsymbol{\mu}}_2, \hat{\nu}_1, \hat{\nu}_2, \hat{\nu}_3, \hat{\nu}_4$ denote the Lagrange multipliers corresponding to $(\hat{\boldsymbol{\alpha}}, \hat{\boldsymbol{\beta}})$, associated with the constraints of (22) in the same order listed here; then the multipliers

$$\begin{aligned} \check{\boldsymbol{\mu}}_1 &= -\check{\mathbf{c}}, \quad \check{\boldsymbol{\mu}}_2 = \hat{\boldsymbol{\mu}}_2 - (\mathbf{B}^T \mathbf{B})^{-1} \mathbf{B}^T \mathbf{K}' \check{\mathbf{c}}, \\ \check{\nu}_1 &= \hat{\nu}_1, \quad \check{\nu}_2 = \hat{\nu}_2, \quad \check{\nu}_3 = \hat{\nu}_3, \quad \check{\nu}_4 = \hat{\nu}_4 \end{aligned} \quad (46)$$

correspond to the solution of (42) given by (45). Recalling that $\hat{\boldsymbol{\theta}} = \hat{\boldsymbol{\mu}}_2$ one readily arrives at (23).

APPENDIX B

EFFICIENT IMPLEMENTATION OF A SPECIAL CASE

In practice, one may effectively ignore the dependencies between the entries $l_m(\mathbf{x})$, $m = 1, \dots, M$, by using a diagonal kernel $K(\mathbf{x}, \mathbf{x}')$ and diagonal basis functions $B_\nu(\mathbf{x})$. Furthermore, one may be interested in modeling all entries identically, as in TPS regression. Then, both the kernel and basis functions become scaled identity matrices; that is, for certain scalar functions $K(\mathbf{x}, \mathbf{x}')$ and $B_\nu(\mathbf{x})$, one has

$$K(\mathbf{x}, \mathbf{x}') := K(\mathbf{x}, \mathbf{x}') \mathbf{I}_M \quad (47)$$

$$B_\nu(\mathbf{x}) := B_\nu(\mathbf{x}) \mathbf{I}_M, \quad \nu = 1, \dots, N_B. \quad (48)$$

Thus, upon defining $\check{\mathbf{K}} \in \mathbb{R}^{N \times N}$ with (n, n') -entry equal to $K(\mathbf{x}_n, \mathbf{x}_{n'})$, and $\check{\mathbf{B}} \in \mathbb{R}^{N \times N_B}$ with (n, ν) -entry equal to $B_\nu(\mathbf{x}_n)$, matrices $\mathbf{K}$ and $\mathbf{B}$ can be written as

$$\mathbf{K} = \check{\mathbf{K}} \otimes \mathbf{I}_M \quad (49)$$

$$\mathbf{B} = \check{\mathbf{B}} \otimes \mathbf{I}_M. \quad (50)$$

Then, the computation of certain matrices involved in the proposed algorithms can be done efficiently as described next.

Start by constructing a selection matrix $\mathbf{S}$ containing ones and zeros such that $\mathbf{A} \odot \mathbf{B} = (\mathbf{A} \otimes \mathbf{B}) \mathbf{S}$, and define $\bar{\boldsymbol{\Phi}}_0 := (\mathbf{I}_N \otimes \mathbf{1}_P^T) \odot \bar{\boldsymbol{\Phi}}$ (c.f. Sec. III-E) and (49) to obtain

$$\begin{aligned} \bar{\boldsymbol{\Phi}}_0^T \mathbf{K} &= [(\mathbf{I}_N \otimes \mathbf{1}_P^T) \odot \bar{\boldsymbol{\Phi}}]^T (\check{\mathbf{K}} \otimes \mathbf{I}_M) \\ &= \mathbf{S}^T [(\mathbf{I}_N \otimes \mathbf{1}_P) \odot \bar{\boldsymbol{\Phi}}^T] (\check{\mathbf{K}} \otimes \mathbf{I}_M) \\ &= \mathbf{S}^T [(\mathbf{I}_N \otimes \mathbf{1}_P) \check{\mathbf{K}} \odot \bar{\boldsymbol{\Phi}}^T] \\ &= \mathbf{S}^T [(\check{\mathbf{K}} \otimes \mathbf{1}_P) \odot \bar{\boldsymbol{\Phi}}^T] \\ &= [(\check{\mathbf{K}} \otimes \mathbf{1}_P^T) \odot \bar{\boldsymbol{\Phi}}]^T. \end{aligned} \quad (51)$$

Likewise, one can verify that

$$\bar{\boldsymbol{\Phi}}_0^T \mathbf{B} = [(\check{\mathbf{B}}^T \otimes \mathbf{1}_P^T) \odot \bar{\boldsymbol{\Phi}}]^T. \quad (52)$$

Using the Kronecker product properties, $\mathbf{P}_B^\perp$ can be written as

$$\begin{aligned} \mathbf{P}_B^\perp &= \mathbf{I}_{MN} - (\check{\mathbf{B}} (\check{\mathbf{B}}^T \check{\mathbf{B}})^{-1} \check{\mathbf{B}}^T) \otimes \mathbf{I}_M \\ &= (\mathbf{I}_N - \check{\mathbf{B}} (\check{\mathbf{B}}^T \check{\mathbf{B}})^{-1} \check{\mathbf{B}}^T) \otimes \mathbf{I}_M = \mathbf{P}_B^\perp \otimes \mathbf{I}_M \end{aligned} \quad (53)$$

where $\mathbf{P}_B^\perp := \mathbf{I}_N - \check{\mathbf{B}} (\check{\mathbf{B}}^T \check{\mathbf{B}})^{-1} \check{\mathbf{B}}^T$. Also, from (30)

$$\begin{aligned} \check{\mathbf{K}} &= \bar{\boldsymbol{\Phi}}_0^T (\mathbf{P}_B^\perp \otimes \mathbf{I}_M) (\check{\mathbf{K}} \otimes \mathbf{I}_M) (\mathbf{P}_B^\perp \otimes \mathbf{I}_M) \bar{\boldsymbol{\Phi}}_0 \\ &= \bar{\boldsymbol{\Phi}}_0^T (\mathbf{P}_B^\perp \check{\mathbf{K}} \mathbf{P}_B^\perp \otimes \mathbf{I}_M) \bar{\boldsymbol{\Phi}}_0 \\ &= ((\mathbf{I}_N \otimes \mathbf{1}_P^T) \odot \bar{\boldsymbol{\Phi}})^T (\mathbf{P}_B^\perp \check{\mathbf{K}} \mathbf{P}_B^\perp \otimes \mathbf{I}_M) \cdot ((\mathbf{I}_N \otimes \mathbf{1}_P^T) \odot \bar{\boldsymbol{\Phi}}) \\ &= \mathbf{S}^T ((\mathbf{I}_N \otimes \mathbf{1}_P^T) \odot \bar{\boldsymbol{\Phi}})^T (\mathbf{P}_B^\perp \check{\mathbf{K}} \mathbf{P}_B^\perp \otimes \mathbf{I}_M) \\ &\quad \cdot ((\mathbf{I}_N \otimes \mathbf{1}_P^T) \odot \bar{\boldsymbol{\Phi}}) \mathbf{S} \\ &= \mathbf{S}^T [(\mathbf{I}_N \otimes \mathbf{1}_P) \mathbf{P}_B^\perp \check{\mathbf{K}} \mathbf{P}_B^\perp (\mathbf{I}_N \otimes \mathbf{1}_P^T) \odot \bar{\boldsymbol{\Phi}}^T \bar{\boldsymbol{\Phi}}] \mathbf{S} \\ &= (\mathbf{I}_N \otimes \mathbf{1}_P) \mathbf{P}_B^\perp \check{\mathbf{K}} \mathbf{P}_B^\perp (\mathbf{I}_N \otimes \mathbf{1}_P^T) \odot \bar{\boldsymbol{\Phi}}^T \bar{\boldsymbol{\Phi}} \\ &= (\mathbf{P}_B^\perp \check{\mathbf{K}} \mathbf{P}_B^\perp \otimes \mathbf{1}_P \mathbf{1}_P^T) \odot \bar{\boldsymbol{\Phi}}^T \bar{\boldsymbol{\Phi}}. \end{aligned} \quad (54)$$

Finally, $\boldsymbol{\theta}$ can be obtained as (cf. Sec. III-E):

$$\boldsymbol{\theta} = \bar{\boldsymbol{\mu}} - [(\check{\mathbf{B}}^T \check{\mathbf{B}})^{-1} \check{\mathbf{B}}^T \check{\mathbf{K}} \otimes \mathbf{I}_M] \hat{\mathbf{c}}. \quad (55)$$

APPENDIX C

PROOF OF THEOREM 1

The following proof extends that of [4, Thm. 1] which cannot accommodate the instantaneous cost from (31) when $\mathcal{L} = \mathcal{L}_{1\epsilon}$ since $\mathcal{L}_{1\epsilon}$ is not differentiable and $\phi(\mathbf{x}_\nu)$ is not constant over ν . Expanding the norms into inner products and employing (32), one finds that

$$\begin{aligned} &\|\mathbf{l}^{(t)} - \mathbf{g}\|_{\mathcal{H}}^2 - \|\mathbf{l}^{(t+1)} - \mathbf{g}\|_{\mathcal{H}}^2 \\ &= -\|\mathbf{l}^{(t+1)} - \mathbf{l}^{(t)}\|_{\mathcal{H}}^2 - 2\langle \mathbf{l}^{(t+1)} - \mathbf{l}^{(t)}, \mathbf{l}^{(t)} - \mathbf{g} \rangle_{\mathcal{H}} \\ &= -\mu_t^2 \|\partial_t \mathcal{C}(\mathbf{l}^{(t)}, \phi(\mathbf{x}_{n_t}), \mathbf{x}_{n_t}, y(\mathbf{x}_{n_t}))\|_{\mathcal{H}}^2 \\ &\quad + 2\langle \mu_t \partial_t \mathcal{C}(\mathbf{l}^{(t)}, \phi(\mathbf{x}_{n_t}), \mathbf{x}_{n_t}, y(\mathbf{x}_{n_t})), \mathbf{l}^{(t)} - \mathbf{g} \rangle_{\mathcal{H}}. \end{aligned} \quad (56)$$

To bound the norm in (56), apply the triangle inequality to (33) to obtain:

$$\begin{aligned} & \|\partial_l \mathcal{C}(\mathbf{l}, \phi(\mathbf{x}_\nu), \mathbf{x}_\nu, y(\mathbf{x}_\nu))\|_{\mathcal{H}} \\ & \leq |\mathcal{L}'_{1\epsilon}(y(\mathbf{x}_\nu) - \phi^T(\mathbf{x}_\nu)\mathbf{l}(\mathbf{x}_\nu))| \|\mathbf{K}(\cdot, \mathbf{x}_\nu)\phi(\mathbf{x}_\nu)\|_{\mathcal{H}} + 2\lambda\|\mathbf{l}\|_{\mathcal{H}}. \end{aligned}$$

Since $\mathcal{L}_{1\epsilon}$ is 1-Lipschitz, $\mathcal{L}'_{1\epsilon}(e) \leq 1 \ \forall e$. Moreover, the reproducing property (cf. Sec. III-A) implies that $\|\mathbf{K}(\cdot, \mathbf{x}_\nu)\phi(\mathbf{x}_\nu)\|_{\mathcal{H}} = (\phi^T(\mathbf{x}_\nu)\mathbf{K}(\mathbf{x}_\nu, \mathbf{x}_\nu)\phi(\mathbf{x}_\nu))^{1/2} < \lambda\|\phi(\mathbf{x}_\nu)\|$, and hence

$$\|\partial_l \mathcal{C}(\mathbf{l}, \phi(\mathbf{x}_\nu), \mathbf{x}_\nu, y(\mathbf{x}_\nu))\|_{\mathcal{H}} \leq \bar{\lambda}\|\phi(\mathbf{x}_\nu)\| + 2\lambda\|\mathbf{l}\|_{\mathcal{H}}. \quad (57)$$

Similarly, from (35) and the triangular inequality, one finds that

$$\|\mathbf{l}^{(t+1)}\|_{\mathcal{H}} \leq (1 - 2\mu_t\lambda)\|\mathbf{l}^{(t)}\|_{\mathcal{H}} + \mu_t\bar{\lambda}\|\phi(\mathbf{x}_\nu)\|.$$

Recalling that $\mathbf{l}^{(1)} = \mathbf{0}$, it is simple to show by induction that $\|\mathbf{l}^{(t)}\|_{\mathcal{H}} \leq U := \bar{\lambda}\bar{\varphi}/2\lambda$ for all t . This fact and $\|\phi(\mathbf{x}_{n_t})\|_2 \leq \bar{\varphi}$ applied to (57) produce

$$\|\partial_l \mathcal{C}(\mathbf{l}, \phi(\mathbf{x}_\nu), \mathbf{x}_\nu, y(\mathbf{x}_\nu))\|_{\mathcal{H}} \leq \bar{\lambda}\bar{\varphi} + 2\lambda U = 2\bar{\lambda}\bar{\varphi}. \quad (58)$$

for all $\mathbf{l}$. On the other hand, the last term in (56) can be bounded by invoking the definition of subgradient:

$$\begin{aligned} & \langle \partial_l \mathcal{C}(\mathbf{l}^{(t)}, \phi(\mathbf{x}_{n_t}), \mathbf{x}_{n_t}, y(\mathbf{x}_{n_t})), \mathbf{l}^{(t)} - \mathbf{g} \rangle_{\mathcal{H}} \\ & \geq \mathcal{C}(\mathbf{l}^{(t)}, \phi(\mathbf{x}_{n_t}), \mathbf{x}_{n_t}, y(\mathbf{x}_{n_t})) - \mathcal{C}(\mathbf{g}, \phi(\mathbf{x}_{n_t}), \mathbf{x}_{n_t}, y(\mathbf{x}_{n_t})). \end{aligned} \quad (59)$$

Combining (58) and (59) with (56) results in

$$\begin{aligned} & \|\mathbf{l}^{(t)} - \mathbf{g}\|_{\mathcal{H}}^2 - \|\mathbf{l}^{(t+1)} - \mathbf{g}\|_{\mathcal{H}}^2 \\ & \geq -4\mu_t^2\bar{\lambda}^2\bar{\varphi}^2 - 2\mu_t[\mathcal{C}(\mathbf{g}, \phi(\mathbf{x}_{n_t}), \mathbf{x}_{n_t}, y(\mathbf{x}_{n_t})) \\ & \quad - \mathcal{C}(\mathbf{l}^{(t)}, \phi(\mathbf{x}_{n_t}), \mathbf{x}_{n_t}, y(\mathbf{x}_{n_t}))]. \end{aligned}$$

Adapting the proof of [4, Prop. 3.1(iii)], it can be shown that if $\mathbf{g} \in \mathcal{H}$ equals the value of $\mathbf{l}$ attaining the infimum on the right hand side of (39), then $\|\mathbf{g}\|_{\mathcal{H}} \leq U$. For such a $\mathbf{g}$ one has

$$\frac{1}{\mu_t}\|\mathbf{l}^{(t)} - \mathbf{g}\|_{\mathcal{H}}^2 - \frac{1}{\mu_{t+1}}\|\mathbf{l}^{(t+1)} - \mathbf{g}\|_{\mathcal{H}}^2 \quad (60)$$

$$= \frac{1}{\mu_t} \left[\|\mathbf{l}^{(t)} - \mathbf{g}\|_{\mathcal{H}}^2 - \|\mathbf{l}^{(t+1)} - \mathbf{g}\|_{\mathcal{H}}^2 \right] \quad (61)$$

$$- \left(\frac{1}{\mu_{t+1}} - \frac{1}{\mu_t} \right) \|\mathbf{l}^{(t+1)} - \mathbf{g}\|_{\mathcal{H}}^2 \quad (62)$$

$$\begin{aligned} & \geq -4\mu_t\bar{\lambda}^2\bar{\varphi}^2 - 2[\mathcal{C}(\mathbf{g}, \phi(\mathbf{x}_{n_t}), \mathbf{x}_{n_t}, y(\mathbf{x}_{n_t})) \\ & \quad - \mathcal{C}(\mathbf{l}^{(t)}, \phi(\mathbf{x}_{n_t}), \mathbf{x}_{n_t}, y(\mathbf{x}_{n_t}))] + \left(\frac{1}{\mu_t} - \frac{1}{\mu_{t+1}} \right) 4U^2 \end{aligned} \quad (63)$$

since $\|\mathbf{l}^{(t+1)} - \mathbf{g}\|_{\mathcal{H}} \leq 2U$. Summing for $t = 1, \dots, T$ and applying $\sum_{t=1}^T \mu_t \leq 2\mu\sqrt{T}$ [24] together with $\sqrt{T+1} - 1 \leq \sqrt{T}$ the result in (39) follows readily.

REFERENCES

- [1] A. Alaya-Feki, S. B. Jemaa, B. Sayrac, P. Houze, and E. Moulines, "Informed spectrum usage in cognitive radio networks: Interference cartography," in *Proc. IEEE Int. Symp. Personal, Indoor Mobile Radio Commun.*, 2008, pp. 1–5.
- [2] A. Argyriou and F. Dinuzzo, "A unifying view of representer theorems," in *Proc. Int. Conf. Mach. Learning*, Beijing, China, Jun. 2014.
- [3] D. D. Ariananda, D. Romero, and G. Leus, "Cooperative compressive power spectrum estimation," in *Proc. of IEEE Sensor Array Multichannel Sig. Process. Workshop*, A Corunha, Spain, Jun. 2014, pp. 97–100.
- [4] J. Audiffren and H. Kadri, "Online learning with multiple operator-valued kernels," Preprint at <http://arXiv.org/abs/1311.0222>, 2013.
- [5] E. Axell, G. Leus, E. G. Larsson, and H. V. Poor, "Spectrum sensing for cognitive radio: State-of-the-art and recent advances," *IEEE Sig. Process. Mag.*, vol. 29, no. 3, pp. 101–116, May 2012.
- [6] J.-A. Bazerque and G. B. Giannakis, "Distributed spectrum sensing for cognitive radio networks by exploiting sparsity," *IEEE Trans. Sig. Process.*, vol. 58, no. 3, pp. 1847–1862, Mar. 2010.
- [7] —, "Nonparametric basis pursuit via kernel-based learning," *IEEE Sig. Process. Mag.*, vol. 28, no. 30, pp. 112–125, 2013.
- [8] J.-A. Bazerque, G. Mateos, and G. B. Giannakis, "Group-lasso on splines for spectrum cartography," *IEEE Trans. Sig. Process.*, vol. 59, no. 10, pp. 4648–4663, Oct. 2011.
- [9] A. Berlinet and C. Thomas-Agnan, *Reproducing Kernel Hilbert Spaces in Probability and Statistics*. Springer, 2004.
- [10] F. L. Bookstein, "Principal warps: Thin-plate splines and the decomposition of deformations," *IEEE Trans. Pattern Anal. Mach. Intel.*, vol. 11, no. 6, pp. 567–585, 1989.
- [11] C. Carmeli, E. De Vito, A. Toigo, and V. Umanita, "Vector valued reproducing kernel Hilbert spaces and universality," *Analysis Appl.*, vol. 8, no. 1, pp. 19–61, 2010.
- [12] E. Dall'Anese, S.-J. Kim, G. B. Giannakis, and S. Pupolin, "Power control for cognitive radio networks under channel uncertainty," *IEEE Trans. Wireless Commun.*, vol. 10, no. 10, pp. 3541–3551, 2011.
- [13] E. Dall'Anese, J.-A. Bazerque, and G. B. Giannakis, "Group sparse lasso for cognitive network sensing robust to model uncertainties and outliers," *Phy. Commun.*, vol. 5, no. 2, pp. 161–172, 2012.
- [14] T. Diethe and M. Girolami, "Online learning with (multiple) kernels: A review," *Neural Comput.*, vol. 25, no. 3, pp. 567–625, Mar. 2013.
- [15] G. Ding, J. Wang, Q. Wu, Y. Yao, F. Song, and T. A. Tsiftsis, "Cellular-base-station-assisted device-to-device communications in TV white space," *IEEE J. Sel. Areas Commun.*, vol. 34, no. 1, pp. 107–121, Jan. 2016.
- [16] M. Gönen and E. Alpaydm, "Multiple kernel learning algorithms," *J. Mach. Learn. Res.*, vol. 12, pp. 2211–2268, 2011.
- [17] M. Gudmundson, "Correlation model for shadow fading in mobile radio systems," *Electron. Letters*, vol. 27, no. 23, pp. 2145–2146, 1991.
- [18] D.-H. Huang, S.-H. Wu, W.-R. Wu, and P.-H. Wang, "Cooperative radio source positioning and power map reconstruction: A sparse Bayesian learning approach," *IEEE Trans. Veh. Technol.*, vol. 64, no. 6, pp. 2318–2332, Jun. 2015.
- [19] B. A. Jayawickrama, E. Dutkiewicz, I. Oppermann, G. Fang, and J. Ding, "Improved performance of spectrum cartography based on compressive sensing in cognitive radio networks," in *Proc. Int. Conf. Commun.*, Budapest, Hungary, Jun. 2013, pp. 5657–5661.
- [20] V. Kecman, M. Vogt, and T. M. Huang, "On the equality of kernel AdaTron and sequential minimal optimization in classification and regression tasks and alike algorithms for kernel machines," in *Proc. European Symp. Artificial Neural Netw.*, Bruges, Belgium, Apr. 2003, pp. 215–222.
- [21] S.-J. Kim, E. Dall'Anese, and G. B. Giannakis, "Cooperative spectrum sensing for cognitive radios using kriged Kalman filtering," *IEEE J. Sel. Topics Sig. Process.*, vol. 5, no. 1, pp. 24–36, Feb. 2011.
- [22] S.-J. Kim and G. B. Giannakis, "Cognitive radio spectrum prediction using dictionary learning," in *Proc. IEEE Global Commun. Conf.*, Atlanta, GA, Dec. 2013.
- [23] S.-J. Kim, N. Jain, G. B. Giannakis, and P. Forero, "Joint link learning and cognitive radio sensing," in *Proc. Asilomar Conf. Sig., Syst., Comput.*, Pacific Grove, CA, Nov. 2011.
- [24] J. Kivinen, A. J. Smola, and R. C. Williamson, "Online learning with kernels," *IEEE Trans. Sig. Process.*, vol. 52, no. 8, pp. 2165–2176, 2004.
- [25] O. Mehanna and N. Sidiropoulos, "Frugal sensing: Wideband power spectrum sensing from few bits," *IEEE Trans. Sig. Process.*, vol. 61, no. 10, pp. 2693–2703, May 2013.
- [26] C. A. Micchelli and M. Pontil, "On learning vector-valued functions," *Neural Comput.*, vol. 17, no. 1, pp. 177–204, 2005.
- [27] J. C. Platt, "Fast training of support vector machines using sequential minimal optimization," in *Advances in Kernel Methods*. Cambridge, MA: MIT Press, 1999, pp. 185–208.
- [28] Z. Quan, S. Cui, H. Poor, and A. Sayed, "Collaborative wideband sensing for cognitive radios," *IEEE Sig. Process. Mag.*, vol. 25, no. 6, pp. 60–73, 2008.

- [29] A. Ribeiro and G. B. Giannakis, "Bandwidth-constrained distributed estimation for wireless sensor networks: Part I/II," *IEEE Trans. Sig. Process.*, vol. 54, no. 3/7, pp. 1131–1143/2784–2796, 2006.
- [30] D. Romero, S.-J. Kim, and G. B. Giannakis, "Stochastic semiparametric regression for spectrum cartography," in *Proc. IEEE Int. Workshop Comput. Advances Multi-Sensor Adaptive Process.*, Cancun, Mexico, Dec. 2015, pp. 513–516.
- [31] D. Romero and G. Leus, "Wideband spectrum sensing from compressed measurements using spectral prior information," *IEEE Trans. Sig. Process.*, vol. 61, no. 24, pp. 6232–6246, 2013.
- [32] B. Schölkopf and A. J. Smola, *Learning with Kernels: Support Vector Machines, Regularization, Optimization, and Beyond*. MIT Press, 2002.
- [33] D. J. Sebal and J. A. Bucklew, "Support vector machine techniques for nonlinear equalization," *IEEE Trans. Sig. Process.*, vol. 48, no. 11, pp. 3217–3226, 2000.
- [34] A. J. Smola and B. Schölkopf, "A tutorial on support vector regression," *Stat. Comput.*, vol. 14, no. 3, pp. 199–222, 2004.
- [35] A. J. Smola, B. Schölkopf, and K.-R. Müller, "The connection between regularization operators and support vector kernels," *Neural Netw.*, vol. 11, no. 4, pp. 637–649, 1998.
- [36] R. Tandra and A. Sahai, "SNR walls for signal detection," *IEEE J. Sel. Topics Sig. Process.*, vol. 2, no. 1, pp. 4–17, 2008.
- [37] J. A. Tropp, J. N. Laska, M. F. Duarte, J. K. Romberg, and R. G. Baraniuk, "Beyond Nyquist: Efficient sampling of sparse bandlimited signals," *IEEE Trans. Inf. Theory*, vol. 56, no. 1, pp. 520–544, Jan. 2010.
- [38] G. Vázquez-Vilar and R. López-Valcarce, "Spectrum sensing exploiting guard bands and weak channels," *IEEE Trans. Sig. Process.*, vol. 59, no. 12, pp. 6045–6057, 2011.
- [39] G. Wahba, *Spline Models for Observational Data*, ser. CBMS-NSF Regional Conference Series in Applied Mathematics. SIAM, 1990, vol. 59.
- [40] S. Wang, W. Zhu, and Z.-P. Liang, "Shape deformation: SVM regression and application to medical image segmentation," in *Proc. Int. Conf. Comput. Vis.*, vol. 2, Vancouver, Canada, 2001, pp. 209–216.
- [41] Z. Wang, K. Crammer, and S. Vucetic, "Breaking the curse of kernelization: Budgeted stochastic gradient descent for large-scale svm training," *J. Machine Learning Research*, vol. 13, no. 1, pp. 3103–3131, 2012.
- [42] T. Yucek and H. Arslan, "A survey of spectrum sensing algorithms for cognitive radio applications," *IEEE Commun. Surveys Tutorials*, vol. 11, no. 1, pp. 116–130, 2009.



D. Romero (M'16) received his M.Sc. and Ph.D. degrees in Signal Theory and Communications from the University of Vigo, Spain, in 2011 and 2015, respectively. From Jul. 2015 to Nov. 2016, he was a post-doctoral researcher with the Digital Technology Center and Department of Electrical and Computer Engineering, University of Minnesota, USA. In Dec. 2017, he joined the Department of Information and Communication Technology, University of Agder, Norway, as an associate professor. His main research interests lie in the areas of signal processing, communications, and machine learning.



Seung-Jun Kim (SM'12) received his B.S. and M.S. degrees from Seoul National University in Seoul, Korea in 1996 and 1998, respectively, and his Ph.D. from the University of California at Santa Barbara in 2005, all in electrical engineering. From 2005 to 2008, he worked for NEC Laboratories America in Princeton, New Jersey, as a Research Staff Member. He was with Digital Technology Center and Department of Electrical and Computer Engineering at the University of Minnesota during 2008–2014, where his final title was Research Associate Professor. In August 2014, he joined Department of Computer Science and Electrical Engineering at the University of Maryland, Baltimore County, as an Assistant Professor. His research interests include statistical signal processing, optimization, and machine learning, with applications to wireless communication and networking, future power systems, and (big) data analytics. He is serving as an Associate Editor for IEEE Signal Processing Letters.



G. B. Giannakis (Fellow'97) received his Diploma in Electrical Engr. from the Ntl. Tech. Univ. of Athens, Greece, 1981. From 1982 to 1986 he was with the Univ. of Southern California (USC), where he received his MSc. in Electrical Engineering, 1983, MSc. in Mathematics, 1986, and Ph.D. in Electrical Engr., 1986. He was with the University of Virginia from 1987 to 1998, and since 1999 he has been a professor with the Univ. of Minnesota, where he holds an Endowed Chair in Wireless Telecommunications, a University of Minnesota McKnight

Presidential Chair in ECE, and serves as director of the Digital Technology Center.

His general interests span the areas of communications, networking and statistical signal processing - subjects on which he has published more than 400 journal papers, 680 conference papers, 25 book chapters, two edited books and two research monographs (h-index 123). Current research focuses on learning from Big Data, wireless cognitive radios, and network science with applications to social, brain, and power networks with renewables. He is the (co-) inventor of 28 patents issued, and the (co-) recipient of 8 best paper awards from the IEEE Signal Processing (SP) and Communications Societies, including the G. Marconi Prize Paper Award in Wireless Communications. He also received Technical Achievement Awards from the SP Society (2000), from EURASIP (2005), a Young Faculty Teaching Award, the G. W. Taylor Award for Distinguished Research from the University of Minnesota, and the IEEE Fourier Technical Field Award (2015). He is a Fellow of EURASIP, and has served the IEEE in a number of posts, including that of a Distinguished Lecturer for the IEEE-SP Society.



Roberto López-Valcarce (S'95-M'01) received the Telecommunication Engineering degree from the University of Vigo, Vigo, Spain, in 1995, and the M.S. and Ph.D. degrees in electrical engineering from the University of Iowa, Iowa City, in 1998 and 2000, respectively. During 1995 he was a Project Engineer with Intelsis. He was a *Ramón y Cajal* Postdoctoral Fellow of the Spanish Ministry of Science and Technology from 2001 to 2006. During that period, he was with the Signal Theory and Communications Department, University of Vigo, where

he currently is an Associate Professor. His main research interests include adaptive signal processing, digital communications, and sensor networks, having coauthored over 50 papers in leading international journals. He holds several patents in collaboration with industry.

Dr. López-Valcarce was a recipient of a 2005 Best Paper Award of the IEEE Signal Processing Society. He served as an Associate Editor of the IEEE TRANSACTIONS ON SIGNAL PROCESSING from 2008 to 2011, and as a member of the IEEE Signal Processing for Communications and Networking Technical Committee from 2011 to 2013.