

Inference of Spatio-Temporal Functions Over Graphs via Multikernel Krige Kalman Filtering

Vassilis N. Ioannidis , *Student Member, IEEE*, Daniel Romero , *Member, IEEE*,
and Georgios B. Giannakis , *Fellow, IEEE*

Abstract—Inference of space-time varying signals on graphs emerges naturally in a plethora of network science related applications. A frequently encountered challenge pertains to reconstructing such dynamic processes, given their values over a subset of vertices and time instants. The present paper develops a graph-aware kernel-based krige Kalman filter that accounts for the spatio-temporal variations, and offers efficient online reconstruction, even for dynamically evolving network topologies. The kernel-based learning framework bypasses the need for statistical information by capitalizing on the smoothness that graph signals exhibit with respect to the underlying graph. To address the challenge of selecting the appropriate kernel, the proposed filter is combined with a multikernel selection module. Such a data-driven method selects a kernel attuned to the signal dynamics on-the-fly within the linear span of a preselected dictionary. The novel multikernel learning algorithm exploits the eigenstructure of Laplacian kernel matrices to reduce computational complexity. Numerical tests with synthetic and real data demonstrate the superior reconstruction performance of the novel approach relative to state-of-the-art alternatives.

Index Terms—Graph signal reconstruction, dynamic models on graphs, krige Kalman filtering, multi-kernel learning.

I. INTRODUCTION

A NUMBER of applications involve data that admit a natural representation in terms of node attributes over social, economic, sensor, communication, and biological networks, to name a few [12], [26]. An inference task that emerges in this context is to predict or extrapolate the attributes of all nodes in the network given the attributes of a subset of them. In a finance network, where nodes correspond to stocks and edges capture dependencies among them, one may be interested in predicting the price of all stocks in the network knowing the price of some. This is of paramount importance in applications where

collecting the attributes of all nodes is prohibitive, as is the case when sampling large-scale graphs, or, when the attribute of interest is of sensitive nature, such as the transmission of HIV in a social network. This task was first formulated as reconstructing a *time-invariant* function on a graph [26], [27].

Follow-up reconstruction approaches leverage the notions of graph bandlimitedness [5], [24], sparsity and overcomplete dictionaries [29], smoothness over the graph [13], [27], all of which can be unified as approximations of nonparametric graph functions drawn from a reproducing kernel Hilbert space (RKHS) [21]; see also [10] for semi-parametric alternatives.

In various applications however, the network connectivity and node attributes change over time. Such is the case in e.g. a finance network, where not only the stock prices change over time, but also their inter-dependencies. Hence, maximizing reconstruction performance for these time-varying signals necessitates judicious modeling of the space-time dynamics, especially when samples are scarce.

Inference of *time-varying* graph functions has been so far pursued mainly for slow variations [9], [15], [30]. Temporal dynamics have been modeled in [18] by assuming that the covariance of the function to be reconstructed is available. On the other hand, spatio-temporal reconstruction of generally dynamic graphs has been approached using an extended graph kernel matrix model with a block tridiagonal structure that lends itself to a computationally tractable iterative solver [19]. However, [19] neither relies on a dynamic model of the function variability, nor it provides a tractable method to learn the “best” kernel that fits the data. Similarly, the algorithm in [11] assumes that an appropriate kernel is known. Furthermore, [18], [11], and [19] do not adapt to changes in the spatio-temporal dynamics of the graph function.

The present paper fills this gap by introducing online estimators for time-varying functions on generally dynamic graphs. Specifically, the contribution is threefold.

- C1. A deterministic model for time-varying network processes is proposed, where spatial dependencies are captured by the topology while spatio-temporal dynamics are described through a graph-aware state-space model.
- C2. Based on this model, an algorithm termed kernel krige Kalman filter (KeKriKF) is developed to obtain function estimates by minimizing a kernel ridge regression (KRR) criterion in an online fashion. The proposed solver generalizes the traditional network krige Kalman filter (KriKF) [17], [18], [31], which relies on a probabilistic

Manuscript received November 23, 2017; revised March 2, 2018 and March 25, 2018; accepted March 26, 2018. Date of publication April 20, 2018; date of current version May 10, 2018. The associate editor coordinating the review of this manuscript and approving it for publication was Prof. Oliver Lezoray. This work was supported by NSF under Grants 1442686, 1500713, and 1508993. This paper was presented in part at the 25th European Signal Processing Conference, Kos island, Greece, Aug.–Sep., 2017. (*Corresponding author: Vassilis N. Ioannidis.*)

V. N. Ioannidis and G. B. Giannakis are with the Department of Electrical and Computer Engineering and the Digital Technology Center, University of Minnesota, Minneapolis, MN 55455 USA (e-mail: ioann006@umn.edu; georgios@umn.edu).

D. Romero is with the Department of ICT, University of Agder, Grimstad 4879, Norway (e-mail: daniel.romero@uia.no).

Color versions of one or more of the figures in this paper are available online at <http://ieeexplore.ieee.org>.

Digital Object Identifier 10.1109/TSP.2018.2827328

model. The novel estimator forgoes with assumptions on data distributions and stationarity, by promoting space-time smoothness through dynamic kernels on graphs.

- C3. To select the most appropriate kernel, a multi-kernel (M)KriKF is developed based on the multi-kernel learning (MKL) framework. This algorithm adaptively selects the kernel that “best” fits the data dynamics within the linear span of a prespecified kernel dictionary. The structure of Laplacian kernels is exploited to reduce complexity down to the order of KeKriKF. This complexity is linear in the number of time samples, which renders KeKriKF and MKriKF appealing for online operation.

The rest of the paper is structured as follows. Section II contains preliminaries and states the problem. Section III introduces the spatio-temporal model and develops the KeKriKF. Section IV endows the KeKriKF with an MKL module to obtain the MKriKF. Finally, numerical experiments and conclusions are presented in Sections V and VI, respectively.

Notation: Scalars are denoted by lowercase, column vectors by bold lowercase, and matrices by bold uppercase letters. Superscripts $^\top$ and † respectively denote transpose and pseudo-inverse; $\mathbf{1}_N$ stands for the $N \times 1$ all-one vector; $\text{diag}\{\mathbf{x}\}$ corresponds to a diagonal matrix with the entries of $\mathbf{x}$ on its diagonal, while $\text{diag}\{\mathbf{X}\}$ is a vector holding the diagonal entries of $\mathbf{X}$; and $\mathcal{N}(\mu, \sigma^2)$ a Gaussian distribution with mean μ and variance σ^2 . Finally, if $\mathbf{A}$ is a matrix and $\mathbf{x}$ a vector, then $\|\mathbf{x}\|_{\mathbf{A}}^2 := \mathbf{x}^\top \mathbf{A}^{-1} \mathbf{x}$ and $\|\mathbf{x}\|_2^2 := \mathbf{x}^\top \mathbf{x}$.

II. PROBLEM STATEMENT AND PRELIMINARIES

Consider a time-varying graph $\mathcal{G}_t := (\mathcal{V}, \mathbf{A}_t)$, $t = 1, 2, \dots$, where $\mathcal{V} := \{v_1, \dots, v_N\}$ denotes the vertex set, and $\mathbf{A}_t$ the $N \times N$ adjacency matrix, whose (n, n') -th entry $A_{n,n'}(t)$ is the nonnegative weight of the edge connecting vertices v_n and $v_{n'}$ at time t . The edge set is $\mathcal{E}_t := \{(v_n, v_{n'}) \in \mathcal{V} \times \mathcal{V} : A_{n,n'}(t) \neq 0\}$, and two vertices v and v' are *connected* at time t if $(v, v') \in \mathcal{E}_t$. The graphs $\{\mathcal{G}_t\}_t$ in this paper are undirected and have no self-loops, which means that $\mathbf{A}_t = \mathbf{A}_t^\top$ and $A_{n,n}(t) = 0$, $\forall t, n$. The Laplacian matrix is $\mathbf{L}_t := \text{diag}\{\mathbf{A}_t \mathbf{1}_N\} - \mathbf{A}_t$, and is positive semidefinite provided that $A_{n,n'}(t) \geq 0$, $\forall n, n', t$.

A time-varying graph function is a map $f : \mathcal{V} \times \mathcal{T} \rightarrow \mathbb{R}$, where $\mathcal{T} := \{1, 2, \dots\}$ is the set of time indices. Specifically, $f(v_n, t)$ represents the value of the attribute of interest at node n and time t , e.g. the closing price of the n -th stock on the t -th day. Vector $\mathbf{f}_t := [f(v_1, t), \dots, f(v_N, t)]^\top \in \mathbb{R}^N$ collects the function values at time t .

Suppose that S_t noisy observations $y(v_{n_s}, t) = f(v_{n_s}, t) + e(v_{n_s}, t)$, $s = 1, \dots, S_t$, are available at time t , where $S_t := \{n_1, \dots, n_{S_t}\}$ contains the sampled indices $1 \leq n_1 \leq \dots \leq n_{S_t} \leq N$, and $e(v_{n_s}, t)$ captures the observation error.

With $\mathbf{y}_t := [y(v_{n_1}, t), \dots, y(v_{n_{S_t}}, t)]^\top$ and $\mathbf{e}_t := [e(v_{n_1}, t), \dots, e(v_{n_{S_t}}, t)]^\top$, the observation model in vector-matrix form is

$$\mathbf{y}_t = \mathbf{S}_t \mathbf{f}_t + \mathbf{e}_t, \quad t = 1, 2, \dots \quad (1)$$

where $\mathbf{S}_t \in \{0, 1\}^{S_t \times N}$ selects the sampled entries of $\mathbf{f}_t$.

Given $\mathbf{y}_\tau$, $\mathbf{S}_\tau$, and $\mathbf{A}_\tau$ for $\tau = 1, \dots, t$, the goal of this paper is to reconstruct $\mathbf{f}_t$ at each t . The estimators should operate in an

TABLE I
EXAMPLES OF LAPLACIAN KERNELS AND THEIR ASSOCIATED SPECTRAL WEIGHT FUNCTIONS

Kernel name	Function	Parameters
Diffusion kernel [13]	$r(\lambda) = \exp\{\sigma^2 \lambda / 2\}$	$\sigma^2 \geq 0$
p -step random walk [27]	$r(\lambda) = (a - \lambda)^{-p}$	$a \geq 2, p$
Regularized Laplacian [26], [27], [32]	$r(\lambda) = 1 + \sigma^2 \lambda$	$\sigma^2 \geq 0$
Bandlimited [21]	$r(\lambda_n) = \begin{cases} 1/\beta & 1 \leq n \leq B \\ \beta & \text{otherwise} \end{cases}$	$\beta > 0, B$
Band-rejection	$r(\lambda_n) = \begin{cases} \beta & k \leq n \leq N-l \\ 1/\beta & \text{otherwise} \end{cases}$	$\beta > 0, k, l$

online fashion, which means that the computational complexity per time slot t must not grow with t . Observe that no statistical information is assumed available in our formulation.

A. Kernel-Based Reconstruction

Aiming ultimately at the time-varying $\mathbf{f}_t$, it is instructive to outline the kernel-based reconstruction of a time-invariant $\mathbf{f} := [f_1, \dots, f_N]$ given $\mathcal{G} := (\mathcal{V}, \mathbf{A})$, and using samples $\mathbf{y} = \mathbf{S}\mathbf{f} + \mathbf{e} \in \mathbb{R}^S$, where $\mathbf{S} \in \{0, 1\}^{S \times N}$ and $S < N$.

Relying on regularized least-squares (LS), we obtain

$$\hat{\mathbf{f}} = \arg \min_{\mathbf{f}} \|\mathbf{y} - \mathbf{S}\mathbf{f}\|_2^2 + \mu g(\mathbf{f}) \quad (2)$$

where $\mu > 0$ and the regularizer $g(\mathbf{f})$ promotes estimates with a certain structure.

For example, the so-called *Laplacian* regularizer

$$g_{\text{LR}}(\mathbf{f}) := (1/2) \sum_{n=1}^N \sum_{n'=1}^N A_{n,n'} (f_n - f_{n'})^2 \quad (3)$$

promotes smooth function estimates with similar values at vertices connected by strong links (large $A_{n,n'}$), since $g_{\text{LR}}(\mathbf{f})$ is small when $\mathbf{f}$ is smooth. It turns out that $g_{\text{LR}}(\mathbf{f}) = \mathbf{f}^\top \mathbf{L} \mathbf{f}$; see e.g. [12, Ch. 2]. For a scalar function $r(\mathbf{L})$ a general graph kernel family of regularizers is obtained as $g_{\text{KR}}(\mathbf{f}) = \mathbf{f}^\top \mathbf{K}^\dagger \mathbf{f} = \|\mathbf{f}\|_{\mathbf{K}}^2$, where the *Laplacian kernel* is defined as

$$\mathbf{K} := r^\dagger(\mathbf{L}) := \mathbf{U}^\top \text{diag}\{r^\dagger(\lambda)\} \mathbf{U} \quad (4)$$

where $\mathbf{U} \in \mathbb{R}^{N \times N}$, and $\lambda \geq 0 \in \mathbb{R}^{N \times 1}$ denote the eigenvector matrix and the eigenvalues of $\mathbf{L}$, since $\mathbf{L} = \mathbf{U} \text{diag}\{\lambda\} \mathbf{U}^\top$. Clearly, $g_{\text{KR}}(\mathbf{f})$ subsumes $g_{\text{LR}}(\mathbf{f})$ for $r(\mathbf{L}) = \mathbf{L}$. Other special cases of $g_{\text{KR}}(\mathbf{f})$ that will be tested in the simulations are collected in Table I, and the scalar functions are plotted in Fig 1. Prior knowledge about the properties of $\mathbf{f}$ may guide the selection of the appropriate $r(\cdot)$. For instance, a diffusion kernel accounts for smoothness of $\mathbf{f}$, as well as the prior that $\mathbf{f}$ is generated by a graph diffusion process. Data-driven selection techniques follow in Section IV.

Further broadening the scope of the generalized Laplacian kernel regularizers, one may set $g(\mathbf{f}) = \|\mathbf{f}\|_{\mathbf{K}}^2$ for an arbitrary positive semidefinite matrix $\mathbf{K}$, not necessarily a Laplacian kernel. These regularizers give rise to the family of *kernel ridge regression* (KRR) estimators

$$\hat{\mathbf{f}} := \arg \min_{\mathbf{f}} \frac{1}{S} \|\mathbf{y} - \mathbf{S}\mathbf{f}\|_2^2 + \mu \|\mathbf{f}\|_{\mathbf{K}}^2 \quad (5)$$

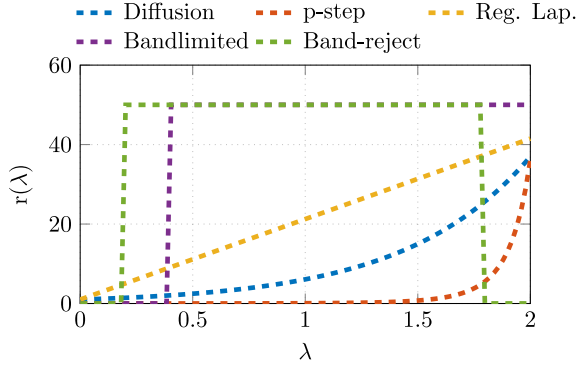


Fig. 1. Laplacian kernels (Diffusion $\sigma = 1.9$, p-step random walk $\alpha = 2.55$, $p = 6$, Regularized Laplacian $\sigma = 4.5$, $\beta = 50$, Bandwidth $B = 20$, $\beta = 50$, Band-reject $k = 10$, $l = 10$).

where $\mu > 0$ controls the effect of the regularizer with respect to the fitting term $S^{-1}\|\mathbf{y} - \mathbf{S}\mathbf{f}\|_2^2$. KRR estimators have well-documented merits and solid grounds on statistical learning theory; see e.g. [23].

So far, signal $\mathbf{f}$ was assumed deterministic. To present a probabilistic interpretation of KRR suppose that $\mathbf{f}$ is zero-mean with $\mathbf{C} := \mathbb{E}[\mathbf{f}\mathbf{f}^\top]$, and that the entries of $\mathbf{e}$ are uncorrelated with each other and with $\mathbf{f}$, and $\sigma_e^2 := S^{-1}\mathbb{E}[\|\mathbf{e}\|_2^2]$. In this setting, the KRR estimator (5) reduces to the linear minimum mean-square error (LMMSE) estimator if $\mu S = \sigma_e^2$ and $\mathbf{K} = \mathbf{C}$. Thus, KRR generalizes LMMSE and can be interpreted as the LMMSE estimator of a random signal $\mathbf{f}$ with covariance matrix $\mathbf{K}$; see [21, Proposition 2].

III. KERNEL KRIGED KALMAN FILTER

This section presents a space-time varying model that is capable of accommodating fairly general forms of spatio-temporal dynamics. Building on this model, a novel online KRR estimator will be subsequently developed for graph functions over time-varying graphs.

A. Spatio-Temporal Model

An immediate approach to estimate $\mathbf{f}_t$ is to apply (5) separately per slot t . This yields the instantaneous estimator

$$\hat{\mathbf{f}}_t^{(\nu)} := \arg \min_{\mathbf{f}} \frac{1}{S_t} \|\mathbf{y}_t - \mathbf{S}_t \mathbf{f}\|_2^2 + \mu \|\mathbf{f}\|_{\mathbf{K}_t}^2 \quad (6)$$

where $\mathbf{K}_t > \mathbf{0}$ is a per-slot preselected kernel matrix, and superscript ν will be explained later. Unfortunately, such an approach does not account for the possible dynamics relating $\mathbf{f}_t$ to $\mathbf{f}_{t-1}$. However, leveraging dependencies across slots can benefit the estimator of $\mathbf{f}_t$ from observations $\{\mathbf{y}_\tau\}_{\tau \neq t}$.

To circumvent the aforementioned limitation, consider modeling the function of interest as

$$f(v_n, t) = f^{(\nu)}(v_n, t) + f^{(x)}(v_n, t) \quad (7)$$

where $f^{(\nu)}$ captures arbitrary (even fast) temporal dynamics across sampling intervals and can be interpreted as an instantaneous component, while $f^{(x)}$ represents a structured (typically slow) varying component.

As an example, consider stock price prediction, where $f^{(\nu)}$ accounts for instantaneous changes caused e.g. by political

statements or company announcements at t relative to $t - 1$, while $f^{(x)}$ captures the steady evolution of the stock market, where stock prices at slot t are closely related to prices of (possibly) other stocks at $t - 1$. Before delving into how these components are modeled, let $\mathbf{f}_t^{(\nu)} := [f^{(\nu)}(v_1, t), \dots, f^{(\nu)}(v_N, t)]^\top$ and $\mathbf{f}_t^{(x)} := [f^{(x)}(v_1, t), \dots, f^{(x)}(v_N, t)]^\top$, and note that (7) can be cast into vector form as

$$\mathbf{f}_t = \mathbf{f}_t^{(\nu)} + \mathbf{f}_t^{(x)}. \quad (8)$$

Vector $\mathbf{f}_t^{(\nu)}$ can be smooth over its entries ($\mathcal{G}_t$), and captures instantaneous dependence among $\{f(v_n, t)\}_{n=1}^N$. This term models the component of the network process that is either uncorrelated over time or it is uncorrelated with respect to the presumed sampling interval. Indeed, one may consider nonlinear models to capture the dynamics of $\mathbf{f}^{(\nu)}$. However, such an approach goes beyond the scope of this submission. On the other hand, $\mathbf{f}_t^{(x)}$ is smooth not only over $\mathcal{G}_t$ but also over time, and models dependencies between $\{f(v_n, t)\}_{n=1}^N$ and their time-lagged versions $\{f(v_n, t - 1)\}_{n=1}^N$.

The smooth evolution of $\mathbf{f}_t^{(x)}$ over time slots adheres to the state equation

$$\mathbf{f}_t^{(x)} = \mathbf{A}_{(t,t-1)} \mathbf{f}_{t-1}^{(x)} + \boldsymbol{\eta}_t, \quad t = 1, 2, \dots \quad (9)$$

where $\mathbf{A}_{(t,t-1)}$ is a graph transition matrix, and $\boldsymbol{\eta}_t := [\eta(v_1, t), \dots, \eta(v_N, t)]^\top \in \mathbb{R}^N$ is termed state noise. Vector $\boldsymbol{\eta}_t$ will be assumed smooth over $\mathcal{G}_t$, meaning $\eta(v_n, t)$ is expected to be similar to $\eta(v_{n'}, t)$ if $A_{n,n'}(t) \neq 0$. The recursion in (9) is the graph counterpart of a vector autoregressive model (VARM) of order one (see e.g. [16], [25]), and will lead to computationally efficient online KRR estimators of $\mathbf{f}_t$ that account for temporal dynamics [25].

Model (8) can be thought of as the graph counterpart of the model adopted in [31] to derive the kriged Kalman filter. In our context here, $\mathbf{f}_t^{(\nu)}$ describes small-scale *spatial fluctuations* within slot t , whereas $\mathbf{f}_t^{(x)}$ captures the so-called *trend* across slots. Furthermore, (8) generalizes the model used in [18], where $\mathbf{A}_{(t,t-1)} = \mathbf{I}_N$, for network delay prediction, where $\mathbf{f}_t^{(\nu)}$ represents the propagation, transmission, and processing delays and $\mathbf{f}_t^{(x)}$ the queuing delay at each router that is affected.

Remark 1: The transition matrix $\mathbf{A}_{(t,t-1)}$ can be interpreted as the $N \times N$ adjacency of a generally directed “transition graph” that relates $\{f^{(x)}(v_n, t - 1)\}_{n=1}^N$ to $\{f^{(x)}(v_n, t)\}_{n=1}^N$. To avoid estimation of $\mathbf{A}_{(t,t-1)}$, the random walk model is motivated, where $\mathbf{A}_{(t,t-1)} = c\mathbf{I}_N$ with $c > 0$. On the other hand, adherence to the graph, prompts the selection $\mathbf{A}_{(t,t-1)} = c\mathbf{A}$, in which case (9) amounts to a diffusion process on a time-invariant $\mathcal{G}$ [24].

B. KeKriKF Algorithm

This section develops an online algorithm to estimate $\mathbf{f}_t$, given (1) and $\{\mathbf{y}_\tau, \mathbf{S}_\tau, \mathbf{A}_\tau, \mathbf{A}_{(\tau,\tau-1)}\}_{\tau=1}^t$ for the spatio-temporal model of $\mathbf{f}_t$ in (8) and (9). Unfortunately, $\{\mathbf{f}_\tau^{(\nu)}\}$ and $\{\mathbf{f}_\tau^{(x)}\}$ cannot be obtained by solving the system of equations comprising (1), (8), and (9) over time even if $\mathbf{e}_\tau = \mathbf{0}$ and $\boldsymbol{\eta}_\tau = \mathbf{0}$, $\forall \tau$; simply because after replacing $\mathbf{f}_\tau$ with

$\mathbf{f}_\tau^{(\chi)} + \mathbf{f}_\tau^{(\nu)} \forall \tau$, the estimation task involves $2Nt$ unknowns, namely $\{\mathbf{f}_\tau^{(\chi)}, \mathbf{f}_\tau^{(\nu)}\}_{\tau=1}^t$, and only $\tilde{S} + Nt$ equations, where $\tilde{S} := \sum_{\tau=1}^t S_\tau$ and $\tilde{S} \leq Nt$. To obtain a solution to this underdetermined problem, one must exploit the model structure. Extending the KRR estimator in (5) to time-varying functions, suppose we wish to

$$\begin{aligned} & \underset{\{\mathbf{f}_\tau^{(\chi)}, \mathbf{f}_\tau^{(\nu)}\}_{\tau=1}^t}{\text{minimize}} \sum_{\tau=1}^t \frac{1}{S_\tau} \|\mathbf{y}_\tau - \mathbf{S}_\tau \mathbf{f}_\tau^{(\chi)} - \mathbf{S}_\tau \mathbf{f}_\tau^{(\nu)}\|^2 \\ & + \mu_1 \sum_{\tau=1}^t \|\mathbf{f}_\tau^{(\chi)} - \mathbf{A}_{(\tau, \tau-1)} \mathbf{f}_{\tau-1}^{(\chi)}\|_{\mathbf{K}_\tau^{(\chi)}}^2 + \mu_2 \sum_{\tau=1}^t \|\mathbf{f}_\tau^{(\nu)}\|_{\mathbf{K}_\tau^{(\nu)}}^2. \end{aligned} \quad (10)$$

where the scalars $\mu_1, \mu_2 \geq 0$ control the trade-off between smoothness and data fit. The first cost is an LS fitting error of the observations. The second cost is a weighted LS error of the slow-varying state (weight matrix promotes spatio-temporal smoothness over $\mathcal{G}_\tau$ and time), and the third term (regularizer) is the weighted ℓ_2 -norm of the fast-varying state (with weight matrix promoting spatial smoothness over $\mathcal{G}_\tau$). When available, prior information about $\{\mathbf{f}_\tau^{(\nu)}, \mathbf{f}_\tau^{(\chi)}\}_{\tau=1}^t$ may steer the selection of suitable kernel matrices; when not available, one can resort to the algorithm in Section IV.

Directly solving (10) per t would not lead to an online algorithm since the complexity of such an approach grows with t ; see Section II. However, we will develop next an efficient *online* algorithm to obtain per slot t estimates $\hat{\mathbf{f}}_{t|t}^{(\chi)}, \hat{\mathbf{f}}_{t|t}^{(\nu)}$ that still account for $\{\mathbf{y}_\tau, \mathbf{S}_\tau, \mathbf{A}_\tau\}_{\tau=1}^t$.

Given $\mathbf{f}_\tau^{(\chi)}$, the first-order necessary conditions for optimality of $\mathbf{f}_\tau^{(\nu)}$ yield [cf. (10)]

$$\mathbf{f}_\tau^{(\nu)} = \mathbf{K}_\tau^{(\nu)} \mathbf{S}_\tau^\top (\bar{\mathbf{K}}_\tau^{(\nu)} + \mu_2 \mathbf{S}_\tau \mathbf{I}_{S_\tau})^{-1} (\mathbf{y}_\tau - \mathbf{S}_\tau \mathbf{f}_\tau^{(\chi)}) \quad (11)$$

where $\bar{\mathbf{K}}_\tau^{(\nu)} := \mathbf{S}_\tau \mathbf{K}_\tau^{(\nu)} \mathbf{S}_\tau^\top$. Notice that the overbar notation indicates $S_\tau \times S_\tau$ matrices or $S_\tau \times 1$ vectors, and recall that without overbar their counterparts have sizes $N \times N$ and $N \times 1$, respectively. Substituting (11) into (10), we arrive at an optimization problem that does not depend on $\mathbf{f}_\tau^{(\nu)}$ for $\tau = 1, \dots, t$. Rewrite next the per slot τ measurement error in (10) using (11) as

$$\begin{aligned} & \frac{1}{S_\tau} \|\mathbf{y}_\tau - \mathbf{S}_\tau \mathbf{f}_\tau^{(\chi)} - \mathbf{S}_\tau \mathbf{f}_\tau^{(\nu)}\|^2 \\ & = \frac{1}{S_\tau} \|\mathbf{y}_\tau - \mathbf{S}_\tau \mathbf{f}_\tau^{(\chi)} - \bar{\mathbf{K}}_\tau^{(\nu)} \\ & \quad \times (\bar{\mathbf{K}}_\tau^{(\nu)} + \mu_2 \mathbf{S}_\tau \mathbf{I}_{S_\tau})^{-1} (\mathbf{y}_\tau - \mathbf{S}_\tau \mathbf{f}_\tau^{(\chi)})\|^2 \\ & = \frac{1}{S_\tau} \|\mathbf{I}_{S_\tau} - \bar{\mathbf{K}}_\tau^{(\nu)} (\bar{\mathbf{K}}_\tau^{(\nu)} + \mu_2 \mathbf{S}_\tau \mathbf{I}_{S_\tau})^{-1}\| \\ & \quad \times (\mathbf{y}_\tau - \mathbf{S}_\tau \mathbf{f}_\tau^{(\chi)})\|^2. \end{aligned} \quad (12a)$$

The matrix inversion lemma asserts for the matrix in square brackets of (12a) that

$$\begin{aligned} & \left[\mathbf{I}_{S_\tau} - \bar{\mathbf{K}}_\tau^{(\nu)} (\bar{\mathbf{K}}_\tau^{(\nu)} + \mu_2 \mathbf{S}_\tau \mathbf{I}_{S_\tau})^{-1} \right] \\ & = \left(\mathbf{I}_{S_\tau} + \frac{1}{\mu_2 S_\tau} \bar{\mathbf{K}}_\tau^{(\nu)} \right)^{-1}. \end{aligned} \quad (12b)$$

Plugging (12b) into (12a) yields

$$\begin{aligned} & = \frac{1}{S_\tau} \left\| \left(\frac{1}{\mu_2 S_\tau} \bar{\mathbf{K}}_\tau^{(\nu)} + \mathbf{I}_{S_\tau} \right)^{-1} (\mathbf{y}_\tau - \mathbf{S}_\tau \mathbf{f}_\tau^{(\chi)}) \right\|^2 \\ & = (\mathbf{y}_\tau - \mathbf{S}_\tau \mathbf{f}_\tau^{(\chi)})^\top \left(\frac{1}{\mu_2} \bar{\mathbf{K}}_\tau^{(\nu)} + \mathbf{S}_\tau \mathbf{I}_{S_\tau} \right)^{-\top} \\ & \quad \times \mathbf{S}_\tau \mathbf{I}_{S_\tau} \left(\frac{1}{\mu_2} \bar{\mathbf{K}}_\tau^{(\nu)} + \mathbf{S}_\tau \mathbf{I}_{S_\tau} \right)^{-1} (\mathbf{y}_\tau - \mathbf{S}_\tau \mathbf{f}_\tau^{(\chi)}). \end{aligned} \quad (12c)$$

Next, we express the regularizer in (10) using (11) for each τ as

$$\begin{aligned} & \mu_2 \|\mathbf{f}_\tau^{(\nu)}\|_{\mathbf{K}_\tau^{(\nu)}}^2 \\ & = (\mathbf{y}_\tau - \mathbf{S}_\tau \mathbf{f}_\tau^{(\chi)})^\top \left(\frac{1}{\mu_2} \bar{\mathbf{K}}_\tau^{(\nu)} + \mathbf{S}_\tau \mathbf{I}_{S_\tau} \right)^{-\top} \\ & \quad \times \frac{1}{\mu_2} \bar{\mathbf{K}}_\tau^{(\nu)} \left(\frac{1}{\mu_2} \bar{\mathbf{K}}_\tau^{(\nu)} + \mathbf{S}_\tau \mathbf{I}_{S_\tau} \right)^{-1} (\mathbf{y}_\tau - \mathbf{S}_\tau \mathbf{f}_\tau^{(\chi)}) \end{aligned} \quad (12d)$$

where the last equality follows from the definition of $\bar{\mathbf{K}}_\tau^{(\nu)}$.

Combining (12c) with (12d) yields

$$\begin{aligned} & \frac{1}{S_\tau} \|\mathbf{y}_\tau - \mathbf{S}_\tau \mathbf{f}_\tau^{(\chi)} - \mathbf{S}_\tau \mathbf{f}_\tau^{(\nu)}\|^2 + \mu_2 \|\mathbf{f}_\tau^{(\nu)}\|_{\mathbf{K}_\tau^{(\nu)}}^2 \\ & = \|\mathbf{y}_\tau - \mathbf{S}_\tau \mathbf{f}_\tau^{(\chi)}\|_{\bar{\mathbf{K}}_\tau^{(\nu)}}^2 \end{aligned} \quad (13)$$

where $\bar{\mathbf{K}}_\tau^{(\nu)} := \frac{1}{\mu_2} \bar{\mathbf{K}}_\tau^{(\nu)} + \mathbf{S}_\tau \mathbf{I}_{S_\tau}$. Using (13) per slot, (10) boils down to

$$\begin{aligned} & \{\hat{\mathbf{f}}_{\tau|t}^{(\chi)}\}_{\tau=1}^t := \arg \min_{\{\mathbf{f}_\tau^{(\chi)}\}_{\tau=1}^t} \sum_{\tau=1}^t \|\mathbf{y}_\tau - \mathbf{S}_\tau \mathbf{f}_\tau^{(\chi)}\|_{\bar{\mathbf{K}}_\tau^{(\nu)}}^2 \\ & + \mu_1 \sum_{\tau=1}^t \|\mathbf{f}_\tau^{(\chi)} - \mathbf{A}_{(\tau, \tau-1)} \mathbf{f}_{\tau-1}^{(\chi)}\|_{\mathbf{K}_\tau^{(\chi)}}^2. \end{aligned} \quad (14)$$

Since (14) is identical to the deterministic formulation of the Kalman filter (KF) applied to a state-space model with state noise covariance $\mathbf{K}_t^{(\chi)}$ and measurement noise covariance $\bar{\mathbf{K}}_t^{(\nu)}$, we deduce that the KF algorithm, see e.g. [28, Ch. 17], applies readily to obtain sequentially the structured per slot t component $\{\hat{\mathbf{f}}_{\tau|t}^{(\chi)}\}_{\tau=1}^t$. After substituting $\{\hat{\mathbf{f}}_{\tau|t}^{(\chi)}\}_{\tau=1}^t$ into (11), we can find also the per slot instantaneous component $\{\hat{\mathbf{f}}_{\tau|t}^{(\nu)}\}_{\tau=1}^t$. The t -th iteration of our so-termed KeKriKF is listed as Algorithm 1.

Summing up, we have established the following result.

Theorem 1: *If $\{\{\hat{\mathbf{f}}_{\tau|t}^{(\chi)}, \hat{\mathbf{f}}_{\tau|t}^{(\nu)}\}_{\tau=1}^t\}_{t=1}^{t_1}$ solves (10) for $t = 1, \dots, t_1$, the KeKriKF iterations summarized in Algorithm 1 for $t = 1, \dots, t_1$ generate the subset of solutions $\{\hat{\mathbf{f}}_{t|t}^{(\chi)}, \hat{\mathbf{f}}_{t|t}^{(\nu)}\}_{t=1}^{t_1}$.*

Clearly, the KeKriKF algorithm comprises two subprocedures: Kalman filtering (steps S1–S6), and kriging (step S7). S3–S5 specify $\mathbf{M}_{t|t-1}$, $\mathbf{M}_{t|t}$, and $\mathbf{G}_t$ that are known in the KF literature as the mean square-error matrices for prediction, correction, and the Kalman gain matrix.

The traditional KriKF has been employed to interpolate stationary processes defined over continuous spatial domains [17],

Algorithm 1: Kernel Krige Kalman Filter (KeKriKF).

Input: $\mathbf{K}_t^{(\chi)}, \mathbf{K}_t^{(\nu)} \in \mathbb{S}_+^N$; $\mathbf{A}_{(t,t-1)} \in \mathbb{R}^{N \times N}$; $\mathbf{y}_t \in \mathbb{R}^{S_t}$;

$$\mathbf{S}_t \in \{0, 1\}^{S_t \times N}; \hat{\mathbf{f}}_{t-1|t-1}^{(\chi)} \in \mathbb{R}^N;$$

$$\mathbf{M}_{t-1|t-1} \in \mathbb{S}_+^N.$$

$$\text{S1. } \tilde{\mathbf{K}}_t^{(\nu)} = \frac{1}{\mu_2} \mathbf{S}_t \mathbf{K}_t^{(\nu)} \mathbf{S}_t^\top + \mathbf{S}_t \mathbf{I}_{S_t}$$

$$\text{S2. } \hat{\mathbf{f}}_{t|t-1}^{(\chi)} = \mathbf{A}_{(t,t-1)} \hat{\mathbf{f}}_{t-1|t-1}^{(\chi)} \quad (\text{prediction})$$

$$\text{S3. } \mathbf{M}_{t|t-1} = \mathbf{A}_{(t,t-1)} \mathbf{M}_{t-1|t-1} \mathbf{A}_{(t,t-1)}^\top + \frac{1}{\mu_1} \mathbf{K}_t^{(\chi)}$$

$$\text{S4. } \mathbf{G}_t = \mathbf{M}_{t|t-1} \mathbf{S}_t^\top (\tilde{\mathbf{K}}_t^{(\nu)} + \mathbf{S}_t \mathbf{M}_{t|t-1} \mathbf{S}_t^\top)^{-1} \quad (\text{gain})$$

$$\text{S5. } \mathbf{M}_{t|t} = (\mathbf{I} - \mathbf{G}_t \mathbf{S}_t) \mathbf{M}_{t|t-1}$$

$$\text{S6. } \hat{\mathbf{f}}_{t|t}^{(\chi)} = \hat{\mathbf{f}}_{t|t-1}^{(\chi)} + \mathbf{G}_t (\mathbf{y}_t - \mathbf{S}_t \hat{\mathbf{f}}_{t|t-1}^{(\chi)}) \quad (\text{correction})$$

$$\text{S7. } \hat{\mathbf{f}}_{t|t}^{(\nu)} = \mathbf{K}_t^{(\nu)} \mathbf{S}_t^\top \tilde{\mathbf{K}}_t^{(\nu)-1} (\mathbf{y}_t - \mathbf{S}_t \hat{\mathbf{f}}_{t|t}^{(\chi)}) \quad (\text{kriging})$$

Output: $\hat{\mathbf{f}}_{t|t}^{(\chi)}; \hat{\mathbf{f}}_{t|t}^{(\nu)}; \mathbf{M}_{t|t}$.

[31], and its derivation follows from a probabilistic linear-minimum mean-square error (LMMSE) criterion that relies on knowledge of second-order statistics [17], [18], [31]. Here, our KeKriKF is derived from a deterministic kernel-based learning framework, which bypasses assumptions on data distributions and stationarity and replaces knowledge of second-order (cross-) covariances with knowledge of $\mathbf{K}_t^{(\nu)}$ and $\mathbf{K}_t^{(\chi)}$. Moreover, different from [7], [15], [18], [30], the novel KeKriKF can accommodate dynamic graph topologies provided $\{\mathbf{K}_t^{(\nu)}, \mathbf{K}_t^{(\chi)}\}_t$ are available.

Remark 2: The complexity of KeKriKF is $\mathcal{O}(N^3)$ per slot. When the underlying graph is large ($N \gg 1$), this complexity can be managed after splitting the graph into N_g sub-graphs each with at most $\lceil N/N_g \rceil$ nodes, and employing consensus-based decentralized KF schemes along the lines of [22].

IV. ONLINE MULTI-KERNEL LEARNING

This section broadens the scope of the KeKriKF algorithm by employing a multi-kernel learning scheme, to bypass the need for selecting an appropriate kernel.

The performance of KRR estimators is well known to heavily depend on the choice of the kernel matrix [21]. Unfortunately, it is difficult to know which kernel matrix is most appropriate for a given problem. To address this issue, an MKL approach is presented that selects a suitable kernel matrix within the linear span of a prespecified dictionary using the available data.

In the following, consider for simplicity that $\mathbf{K}_t^{(\nu)} = \mathbf{K}^{(\nu)}$, $\mathbf{K}_t^{(\chi)} = \mathbf{K}^{(\chi)}$, and $\mathbf{S}_t = \mathbf{S}$, $\forall t$. The kernels in the dictionaries $\mathcal{D}^{(\nu)} := \{\mathbf{K}^{(\nu)}[m] \in \mathbb{S}_+^N\}_{m=1}^{M_\nu}$, and $\mathcal{D}^{(\chi)} := \{\mathbf{K}^{(\chi)}[m] \in \mathbb{S}_+^N\}_{m=1}^{M_\chi}$ will be combined to generate $\mathbf{K}^{(\nu)} = \mathbf{K}^{(\nu)}$ ($\boldsymbol{\theta}^{(\nu)} := \sum_{m=1}^{M_\nu} \theta^{(\nu)}[m] \mathbf{K}^{(\nu)}[m]$) and $\mathbf{K}^{(\chi)} = \mathbf{K}^{(\chi)}(\boldsymbol{\theta}^{(\chi)}) := \sum_{m=1}^{M_\chi} \theta^{(\chi)}[m] \mathbf{K}^{(\chi)}[m]$, where $\boldsymbol{\theta}^{(\nu)} := [\theta^{(\nu)}[1], \dots, \theta^{(\nu)}[M_\nu]]^\top$, $\boldsymbol{\theta}^{(\chi)} := [\theta^{(\chi)}[1], \dots, \theta^{(\chi)}[M_\chi]]^\top \succeq \mathbf{0}$ are coefficients to be determined.

Next, consider expanding the optimization in (10) to obtain $\boldsymbol{\theta}^{(\nu)}, \boldsymbol{\theta}^{(\chi)}$ along with $\{\mathbf{f}_\tau^{(\chi)}, \mathbf{f}_\tau^{(\nu)}\}_{\tau=1}^t$, as follows

$$\begin{aligned} & \underset{\substack{\{\mathbf{f}_\tau^{(\chi)}, \mathbf{f}_\tau^{(\nu)}\}_{\tau=1}^t, \\ \boldsymbol{\theta}^{(\chi)} \succeq \mathbf{0}, \boldsymbol{\theta}^{(\nu)} \succeq \mathbf{0}}}{\text{minimize}} \quad \frac{1}{t} \sum_{\tau=1}^t \frac{1}{S} \|\mathbf{y}_\tau - \mathbf{S} \mathbf{f}_\tau^{(\chi)} - \mathbf{S} \mathbf{f}_\tau^{(\nu)}\|^2 \\ & + \frac{\mu_1}{t} \sum_{\tau=1}^t \|\mathbf{f}_\tau^{(\chi)} - \mathbf{A}_{(\tau,\tau-1)} \mathbf{f}_{\tau-1}^{(\chi)}\|_{\mathbf{K}^{(\chi)}(\boldsymbol{\theta}^{(\chi)})}^2 \\ & + \frac{\mu_2}{t} \sum_{\tau=1}^t \|\mathbf{f}_\tau^{(\nu)}\|_{\mathbf{K}^{(\nu)}(\boldsymbol{\theta}^{(\nu)})}^2 + \rho_\nu \|\boldsymbol{\theta}^{(\nu)}\|_2^2 + \rho_\chi \|\boldsymbol{\theta}^{(\chi)}\|_2^2 \quad (15) \end{aligned}$$

where $\rho_\nu, \rho_\chi \geq 0$ are regularization parameters. The solution to (15) for each t will be denoted as $\{\hat{\mathbf{f}}_{\tau|t}^{(\chi)}, \hat{\mathbf{f}}_{\tau|t}^{(\nu)}\}_{\tau=1}^t \cup \{\hat{\boldsymbol{\theta}}_t^{(\chi)}, \hat{\boldsymbol{\theta}}_t^{(\nu)}\}$. Here, the data-dependent $\{\hat{\boldsymbol{\theta}}_t^{(\chi)}, \hat{\boldsymbol{\theta}}_t^{(\nu)}\}$ select the kernel matrices that “best” capture the data dynamics.

Due to the presence of the weighted norms, namely $\{\|\mathbf{f}_\tau^{(\chi)} - \mathbf{A}_{(\tau,\tau-1)} \mathbf{f}_{\tau-1}^{(\chi)}\|_{\mathbf{K}^{(\chi)}(\boldsymbol{\theta}^{(\chi)})}^2\}_{\tau=1}^t$ and $\{\|\mathbf{f}_\tau^{(\nu)}\|_{\mathbf{K}^{(\nu)}(\boldsymbol{\theta}^{(\nu)})}^2\}_{\tau=1}^t$, the problem in (15) is non-convex. Fortunately, (15) is separately convex in $\{\mathbf{f}_\tau^{(\chi)}, \mathbf{f}_\tau^{(\nu)}\}_{\tau=1}^t, \boldsymbol{\theta}^{(\nu)}, \boldsymbol{\theta}^{(\chi)}$, which motivates the use of alternating minimization (AM) strategies. AM algorithms minimize the objective with respect to every block of variables, while keeping the other variables fixed [8]. Conveniently, if $\boldsymbol{\theta}^{(\nu)}, \boldsymbol{\theta}^{(\chi)}$ are fixed, then (15) reduces to (10), which can be solved by Algorithm 1 for $\hat{\mathbf{f}}_{t|t}^{(\nu)}, \hat{\mathbf{f}}_{t|t}^{(\chi)}$ per slot t ; see Theorem 1. Conversely, $\hat{\boldsymbol{\theta}}_t^{(\chi)}, \hat{\boldsymbol{\theta}}_t^{(\nu)}$ can be obtained for fixed $\{\mathbf{f}_\tau^{(\nu)}, \mathbf{f}_\tau^{(\chi)}\}_{\tau=1}^t$ as specified next.

Theorem 2: Consider minimizing (15) with respect to $\boldsymbol{\theta}^{(\chi)}$ and $\boldsymbol{\theta}^{(\nu)}$ for fixed $\mathbf{f}_\tau^{(\chi)} = \hat{\mathbf{f}}_{\tau|\tau}^{(\chi)}$ and $\mathbf{f}_\tau^{(\nu)} = \hat{\mathbf{f}}_{\tau|\tau}^{(\nu)}$, $\tau = 1, \dots, t$, where $\{\hat{\mathbf{f}}_{\tau|\tau}^{(\chi)}, \hat{\mathbf{f}}_{\tau|\tau}^{(\nu)}\}_{\tau=1}^t$ are given and not necessarily the global minimizers of (15) with respect to $\{\mathbf{f}_\tau^{(\chi)}, \mathbf{f}_\tau^{(\nu)}\}_{\tau=1}^t$. Let $\tilde{\mathbf{f}}_{\tau|\tau}^{(\chi)} := \hat{\mathbf{f}}_{\tau|\tau}^{(\chi)} - \mathbf{A}_{(\tau,\tau-1)} \hat{\mathbf{f}}_{\tau-1|\tau-1}^{(\chi)}$, $\tau = 2, \dots, t$, as well as $\mathbf{R}_t^{(\nu)} = \frac{1}{t} \sum_{\tau=1}^t \hat{\mathbf{f}}_{\tau|\tau}^{(\nu)} \hat{\mathbf{f}}_{\tau|\tau}^{(\nu)\top}$ and $\mathbf{R}_t^{(\chi)} = \frac{1}{t} \sum_{\tau=1}^t \tilde{\mathbf{f}}_{\tau|\tau}^{(\chi)} \tilde{\mathbf{f}}_{\tau|\tau}^{(\chi)\top}$. Then, the minimizers of (15) with respect to $\boldsymbol{\theta}^{(\nu)}$ and $\boldsymbol{\theta}^{(\chi)}$ are

$$\hat{\boldsymbol{\theta}}_t^{(\nu)} = \arg \min_{\boldsymbol{\theta}^{(\nu)} \succeq \mathbf{0}} \text{Tr}\{\mathbf{R}_t^{(\nu)} \mathbf{K}^{(\nu)-1}(\boldsymbol{\theta}^{(\nu)})\} + \frac{\rho_\nu}{\mu_2} \|\boldsymbol{\theta}^{(\nu)}\|_2^2 \quad (16a)$$

$$\hat{\boldsymbol{\theta}}_t^{(\chi)} = \arg \min_{\boldsymbol{\theta}^{(\chi)} \succeq \mathbf{0}} \text{Tr}\{\mathbf{R}_t^{(\chi)} \mathbf{K}^{(\chi)-1}(\boldsymbol{\theta}^{(\chi)})\} + \frac{\rho_\chi}{\mu_1} \|\boldsymbol{\theta}^{(\chi)}\|_2^2. \quad (16b)$$

Proof: To prove (16a), keep in (15) only those terms that depend on $\boldsymbol{\theta}^{(\nu)}$, and replace $\{\mathbf{f}_\tau^{(\nu)}\}_{\tau=1}^t$ with $\{\hat{\mathbf{f}}_{\tau|\tau}^{(\nu)}\}_{\tau=1}^t$. Then, the objective in (15) reduces to $(1/t) \sum_{\tau=1}^t \hat{\mathbf{f}}_{\tau|\tau}^{(\nu)\top} \mathbf{K}^{(\nu)-1}(\boldsymbol{\theta}^{(\nu)}) \hat{\mathbf{f}}_{\tau|\tau}^{(\nu)} + (\rho_\nu/\mu_2) \|\boldsymbol{\theta}^{(\nu)}\|_2^2$. Next, using the linearity and cyclic invariance of the trace it follows that $\text{Tr}\{(1/t) \sum_{\tau=1}^t \hat{\mathbf{f}}_{\tau|\tau}^{(\nu)\top} \mathbf{K}^{(\nu)-1}(\boldsymbol{\theta}^{(\nu)}) \hat{\mathbf{f}}_{\tau|\tau}^{(\nu)}\} = \text{Tr}\{(1/t) \sum_{\tau=1}^t \hat{\mathbf{f}}_{\tau|\tau}^{(\nu)} \hat{\mathbf{f}}_{\tau|\tau}^{(\nu)\top} \mathbf{K}^{(\nu)-1}(\boldsymbol{\theta}^{(\nu)})\} = \text{Tr}\{\mathbf{R}_t^{(\nu)} \mathbf{K}^{(\nu)-1}(\boldsymbol{\theta}^{(\nu)})\}$, which proves (16a). The proof of (16b) follows along the same lines. ■

Algorithm 2: Multi-kernel KriKF (MKriKF).

Input: $\mathcal{D}^{(\nu)}; \mathcal{D}^{(\chi)}; \mathbf{L} = \mathbf{U}^\top \text{diag}\{\boldsymbol{\lambda}\} \mathbf{U}$.

- 1: Initialize: $\hat{\boldsymbol{\theta}}_0^{(\nu)} = \hat{\boldsymbol{\theta}}_0^{(\chi)} = [1, 0, \dots, 0]$, $\hat{\mathbf{f}}_{0|0}^{(\chi)} = \mathbf{0}$,
 $\mathbf{M}_{0|0} = \frac{1}{\mu_1} \mathbf{K}^{(\chi)}[1]$,
 $\boldsymbol{\lambda}^{(\nu)}[m] := \text{diag}\{\mathbf{U} \mathbf{K}^{(\nu)}[m] \mathbf{U}^\top\} \forall m$,
 $\boldsymbol{\lambda}^{(\chi)}[m] := \text{diag}\{\mathbf{U} \mathbf{K}^{(\chi)}[m] \mathbf{U}^\top\} \forall m$.
- 2: **for** $t = 1, 2, \dots$ **do**
- 3: **Input:** $\mathbf{A}_{(t,t-1)} \in \mathbb{R}^{N \times N}; \mathbf{y}_t \in \mathbb{R}^{S_t}; \mathbf{S}_t \in \{0, 1\}^{S_t \times N}$.
- 4: $\mathbf{K}_t^{(\nu)} = \mathbf{K}^{(\nu)}(\hat{\boldsymbol{\theta}}_t^{(\nu)})$
- 5: $\mathbf{K}_t^{(\chi)} = \mathbf{K}^{(\chi)}(\hat{\boldsymbol{\theta}}_t^{(\chi)})$
- 6: $\{\hat{\mathbf{f}}_{t|t}^{(\nu)}, \hat{\mathbf{f}}_{t|t}^{(\chi)}\} = \text{KeKriKF}(\mathbf{K}_{t-1}^{(\chi)}, \mathbf{K}_{t-1}^{(\nu)}, \mathbf{A}_{(t,t-1)},$
 $\mathbf{y}_t, \mathbf{S}_t, \hat{\mathbf{f}}_{t-1|t-1}^{(\chi)}, \mathbf{M}_{t-1|t-1})$
- 7: Update $\mathbf{R}_t^{(\nu)}$ and $\mathbf{R}_t^{(\chi)}$
- 8: $\mathbf{T}_t^{(\nu)} = \mathbf{U}^\top \mathbf{R}_t^{(\nu)} \mathbf{U}$
- 9: $\mathbf{T}_t^{(\chi)} = \mathbf{U}^\top \mathbf{R}_t^{(\chi)} \mathbf{U}$
- 10: $\hat{\boldsymbol{\theta}}_t^{(\nu)} = \text{OKM}(\{\boldsymbol{\lambda}^{(\nu)}[m]\}_{m=1}^{M_\nu}, \mathbf{T}_t^{(\nu)}, \hat{\boldsymbol{\theta}}_{t-1}^{(\nu)})$
- 11: $\hat{\boldsymbol{\theta}}_t^{(\chi)} = \text{OKM}(\{\boldsymbol{\lambda}^{(\chi)}[m]\}_{m=1}^{M_\chi}, \mathbf{T}_t^{(\chi)}, \hat{\boldsymbol{\theta}}_{t-1}^{(\chi)})$
- 12: **Output:** $\hat{\mathbf{f}}_{t|t}^{(\chi)}; \hat{\mathbf{f}}_{t|t}^{(\nu)}; \mathbf{M}_{t|t}$.
- 13: **end for**

Thus, Theorem 2 simplifies the objective that has to be minimized to find $\hat{\boldsymbol{\theta}}_t^{(\chi)}$ and $\hat{\boldsymbol{\theta}}_t^{(\nu)}$. With $\mathbf{K}(\boldsymbol{\theta}) = \sum_{m=1}^M \theta[m] \mathbf{K}[m]$, problems (16a) and (16b) are of the form

$$\hat{\boldsymbol{\theta}} = \arg \min_{\boldsymbol{\theta} \geq \mathbf{0}} \text{Tr}\{\mathbf{R} \mathbf{K}^{-1}(\boldsymbol{\theta})\} + \rho \|\boldsymbol{\theta}\|_2^2 \quad (17)$$

for some $\mathbf{R} \in \mathbb{R}^{N \times N}$, $\rho \geq 0$, and $\mathcal{D} = \{\mathbf{K}[m]\}_{m=1}^M$. Due to their resemblance to *covariance matching* [20], problem (17), and hence (16a) and (16b) will be referred to as *kernel matching*.

Theorem 2 suggests an online AM procedure to approximate the solution to (15), where Algorithm 1 and a solver for (17) termed online kernel matching (OKM) are executed alternately. This is summarized as Algorithm 2, and it is termed multi-kernel KriKF (MKriKF). Algorithm 2 does not generally find a global optimum of (15); yet, finding such an optimum may not be critical in practice, since it cannot be computed in polynomial time.

The rest of this section develops the OKM algorithm for solving (17) when $\mathcal{D}$ comprises Laplacian kernels. The first step is to exploit the fact that all Laplacian kernel matrices associated with a given graph have common eigenvectors.

Proposition 1: Consider the eigenvalue decompositions $\{\mathbf{K}[m] = \mathbf{U} \text{diag}\{\boldsymbol{\lambda}[m]\} \mathbf{U}^\top\}_{m=1}^M$ and let $\mathbf{T} := \mathbf{U}^\top \mathbf{R} \mathbf{U}$. Upon defining $\boldsymbol{\Lambda}(\boldsymbol{\theta}) := \text{diag}\{\sum_{m=1}^M \theta[m] \boldsymbol{\lambda}[m]\}$ and $\phi(\boldsymbol{\theta}) := \text{Tr}(\mathbf{T} \boldsymbol{\Lambda}^{-1}(\boldsymbol{\theta})) + \rho \|\boldsymbol{\theta}\|_2^2$, (17) can be equivalently written as

$$\hat{\boldsymbol{\theta}} = \arg \min_{\boldsymbol{\theta} \geq \mathbf{0}} \phi(\boldsymbol{\theta}) \quad (18)$$

Proof: Since $\mathbf{K}(\boldsymbol{\theta}) = \sum_{m=1}^M \theta[m] \mathbf{U} \text{diag}\{\boldsymbol{\lambda}[m]\} \mathbf{U}^\top = \mathbf{U} \boldsymbol{\Lambda}(\boldsymbol{\theta}) \mathbf{U}^\top$, (18) follows by noting that $\text{Tr}\{\mathbf{R} \mathbf{K}^{-1}(\boldsymbol{\theta})\} = \text{Tr}\{\mathbf{R} \mathbf{U} \boldsymbol{\Lambda}^{-1}(\boldsymbol{\theta}) \mathbf{U}^\top\} = \text{Tr}\{\mathbf{U}^\top \mathbf{R} \mathbf{U} \boldsymbol{\Lambda}^{-1}(\boldsymbol{\theta})\} = \text{Tr}\{\mathbf{T} \boldsymbol{\Lambda}^{-1}(\boldsymbol{\theta})\}$. ■

Algorithm 3: Online Kernel Matching (OKM).

Input: $\{\boldsymbol{\lambda}[m]\}_{m=1}^M; \mathbf{T}_t \in \mathbb{S}_+^N; \hat{\boldsymbol{\theta}}_{t-1} \in \mathbb{R}_+^M$.

- 1: Initialize: $\boldsymbol{\theta}^0 = \hat{\boldsymbol{\theta}}_{t-1}$,
- 2: **while** stopping_criterion not met **do**
- 3: $\boldsymbol{\theta}^{k+1} = [\boldsymbol{\theta}^k - s^k \nabla \phi(\boldsymbol{\theta}^k)]^+$
- 4: $k \leftarrow k + 1$
- 5: **end while**

Output: $\hat{\boldsymbol{\theta}}_t$.

Proposition 1 establishes that (17) can be expressed as (18) when the kernels in $\mathcal{D}$ share eigenvectors, as is the case of Laplacian kernels; cf. Section II-A.

Proposition 2: When $\boldsymbol{\theta} \succeq \mathbf{0}$, function $\phi(\boldsymbol{\theta})$ is strongly convex and differentiable with gradient

$$\nabla \phi(\boldsymbol{\theta}) = \mathbf{v}(\boldsymbol{\theta}) + 2\rho \boldsymbol{\theta} \quad (19)$$

where $\mathbf{v}(\boldsymbol{\theta}) := -[\text{Tr}\{\text{diag}\{\tilde{\boldsymbol{\lambda}}[1]\} \mathbf{T}\}, \dots, \text{Tr}\{\text{diag}\{\tilde{\boldsymbol{\lambda}}[M]\} \mathbf{T}\}]$, with $\tilde{\boldsymbol{\lambda}}[m] := [\tilde{\lambda}_1[m], \dots, \tilde{\lambda}_N[m]]^\top$ and $\tilde{\lambda}_n[m] := \lambda_n[m] / (\sum_{\mu=1}^M \theta[\mu] \lambda_n[\mu])^2$.

Proof: Because $\mathbf{T}$ is a positive semidefinite matrix and $\boldsymbol{\lambda}[m] \succeq \mathbf{0} \forall m$, it can be easily seen that $\text{Tr}\{\mathbf{T} \boldsymbol{\Lambda}^{-1}(\boldsymbol{\theta})\}$ is convex over $\boldsymbol{\theta} \succeq \mathbf{0}$. And since $\rho \|\boldsymbol{\theta}\|_2^2$ is strongly convex, it follows by its definition that $\phi(\boldsymbol{\theta})$ is strongly convex. To obtain the gradient observe that

$$\frac{\partial \phi}{\partial \theta[m]} = -\text{Tr}\{\boldsymbol{\Lambda}^{-1}(\boldsymbol{\theta}) \text{diag}\{\boldsymbol{\lambda}[m]\} \boldsymbol{\Lambda}^{-1}(\boldsymbol{\theta}) \mathbf{T}\} + 2\rho \theta[m] \quad (20)$$

and $\boldsymbol{\Lambda}^{-1}(\boldsymbol{\theta}) \text{diag}\{\boldsymbol{\lambda}[m]\} \boldsymbol{\Lambda}^{-1}(\boldsymbol{\theta}) = \text{diag}\{\tilde{\boldsymbol{\lambda}}[m]\}$. ■

As (18) entails a strongly convex and differentiable objective, and projections on its feasible set are easy to obtain, we are motivated to solve (18) through projected gradient descent (PGD) [6]. Besides its simplicity, PGD converges linearly to the global minimum of (18). The general PGD iteration is

$$\boldsymbol{\theta}^{k+1} = [\boldsymbol{\theta}^k - s^k \nabla \phi(\boldsymbol{\theta}^k)]^+, \quad k = 0, 1, \dots \quad (21)$$

where s^k is the stepsize chosen e.g. by the Armijo rule [6], $\boldsymbol{\theta}^0$ is a feasible initial step, and $[\cdot]^+$ denotes projection on the non-negative orthant $\{\boldsymbol{\theta} : \theta[m] \geq 0, m = 1, \dots, M\}$. The overall algorithm is termed OKM, and it is listed as Algorithm 3.

Observe that $\boldsymbol{\theta}^0$ in Algorithm 3 is initialized with its output in the previous iterate, namely $\hat{\boldsymbol{\theta}}_{t-1}$. This is a warm start that considerably speeds up convergence of Algorithm 3 since $\phi(\boldsymbol{\theta})$ is expected to change slowly across the iterations in Algorithm 2. An interesting byproduct of the OKM algorithm is its ability to adapt to changes in the spatio-temporal dynamics of the graph functions by adjusting the coefficients $\{\hat{\boldsymbol{\theta}}_t^{(\nu)}, \hat{\boldsymbol{\theta}}_t^{(\chi)}\}_t$, and consequently the kernel matrices.

In view of Proposition 2, finding each entry of $\nabla \phi(\boldsymbol{\theta})$ in Algorithm 3 requires $\mathcal{O}(N)$ operations. Computing the gradient through (19) exploits the common eigenvectors of $\{\mathbf{K}[m]\}_{m=1}^M$, and avoids the inversion of the $N \times N$ matrix $\mathbf{K}(\boldsymbol{\theta})$ that is required when calculating the gradient for the general formulation (17), where $\{\mathbf{K}[m]\}_{m=1}^M$ need not share eigenvectors.

The complexity of evaluating the gradient is therefore reduced from a prohibitive $\mathcal{O}(N^3M)$ for general kernels to an affordable $\mathcal{O}(NM)$ for Laplacian kernels, which amounts to considerable computational savings especially for large-scale networks. With K denoting the number of PGD iterations for convergence, the overall computational complexity of OKM is therefore $\mathcal{O}(NMK)$. Typically, $N^3 \geq NMK$ and hence the complexity of Algorithm 2 is $\mathcal{O}(N^3)$, while learning the appropriate linear combination of kernels through MKL does not increase the complexity order that can be further reduced as suggested in Remark 2.

Selecting the dictionary and its size clearly depends on the amount of prior information available, and the complexity that can be afforded by the MKL optimization that follows up. Desirable attributes such as smoothness, bandlimitedness, and diffusion effects can prompt inclusion of corresponding kernels over a grid of their parameters - the case present in our simulation tests.

Remark 3: The algorithms in this section adopted a fixed kernel dictionary over time, namely $\mathcal{D} = \{\mathbf{K}[m] \in \mathbb{S}_+^N\}_{m=1}^M$. If the topology changes over time, the Laplacian kernel matrices change as well, cf. (4). To accommodate this scenario, one can restart Algorithm 2 whenever the topology changes, say at time t_c , and initialize $\hat{\mathbf{f}}_{0|0}^{(\chi)} \leftarrow \hat{\mathbf{f}}_{t_c|t_c}^{(\chi)}$, $\mathbf{M}_{0|0} \leftarrow \mathbf{M}_{t_c|t_c}$, as well as replace the Laplacian kernels in $\mathcal{D}$ with the ones corresponding to the new topology.

Remark 4: To accommodate a certain degree of nonstationarity one may consider using the following matrices

$$\tilde{\mathbf{R}}_t^{(\nu)} = \sum_{\tau=1}^t \gamma_\nu^{t-\tau} \hat{\mathbf{f}}_{\tau|\tau}^{(\nu)} \hat{\mathbf{f}}_{\tau|\tau}^{(\nu)\top} + \gamma_\nu^t \mathbf{I} \quad (22a)$$

$$\tilde{\mathbf{R}}_t^{(\chi)} = \sum_{\tau=1}^t \gamma_\chi^{t-\tau} \tilde{\mathbf{f}}_{\tau|\tau}^{(\chi)} \tilde{\mathbf{f}}_{\tau|\tau}^{(\chi)\top} + \gamma_\chi^t \mathbf{I} \quad (22b)$$

instead of $\mathbf{R}_t^{(\nu)}$ and $\mathbf{R}_t^{(\chi)}$, where $\gamma_\chi, \gamma_\nu \in (0, 1)$ are forgetting factors that weigh exponentially past observations, and ensure invertibility of matrices $\tilde{\mathbf{R}}_t^{(\nu)}$ and $\tilde{\mathbf{R}}_t^{(\chi)}$. Moreover, $\tilde{\mathbf{R}}_t^{(\nu)}$ and $\tilde{\mathbf{R}}_t^{(\chi)}$ can be updated recursively as

$$\tilde{\mathbf{R}}_t^{(\nu)} = \gamma_\nu \tilde{\mathbf{R}}_{t-1}^{(\nu)} + \hat{\mathbf{f}}_{t|t}^{(\nu)} \hat{\mathbf{f}}_{t|t}^{(\nu)\top} \quad (23a)$$

$$\tilde{\mathbf{R}}_t^{(\chi)} = \gamma_\chi \tilde{\mathbf{R}}_{t-1}^{(\chi)} + \tilde{\mathbf{f}}_{t|t}^{(\chi)} \tilde{\mathbf{f}}_{t|t}^{(\chi)\top} \quad (23b)$$

which significantly reduces the required memory for the computation with respect to (22), since $\{\hat{\mathbf{f}}_{\tau|\tau}^{(\nu)}, \tilde{\mathbf{f}}_{\tau|\tau}^{(\chi)}\}_{\tau=1}^{t-1}$ need not be stored.

Remark 5: The algorithms presented in this paper can be generalized to account for a VARMA of order L , $\mathbf{f}_t^{(\chi)} = \sum_{l=1}^L \mathbf{A}_{(t,t-l)} \mathbf{f}_{t-l}^{(\chi)} + \boldsymbol{\eta}_t$. Towards that end, consider the $LN \times 1$ extended state vector $\tilde{\mathbf{f}}_t^{(\chi)} := [(\tilde{\mathbf{f}}_t^{(\chi)})^\top, \dots, (\tilde{\mathbf{f}}_{t-L+1}^{(\chi)})^\top]^\top$, the $S_t \times LN$ matrix $\tilde{\mathbf{S}}_t := [\mathbf{S}_t, \mathbf{0}, \dots, \mathbf{0}]$, the $LN \times LN$ matrices $\tilde{\mathbf{A}}_{(t,t-1)}$ with block entries $\{[\tilde{\mathbf{A}}_{(t,t-1)}]_{1,l} = \mathbf{A}_{(t,t-l)}\}_{l=1}^L$, $\{[\tilde{\mathbf{A}}_{(t,t-1)}]_{l,l} = \mathbf{I}_N\}_{l=2}^L$, and the rest zero, and $\tilde{\mathbf{K}}_t^{(\chi)}$ with block entries $[\tilde{\mathbf{K}}_t^{(\chi)}]_{1,1} = \mathbf{K}_t^{(\chi)}$, $\{[\tilde{\mathbf{K}}_t^{(\chi)}]_{l,l} = \mathbf{I}_N\}_{l=2}^L$, and

the rest zero. The KeKriKF algorithm can then be readily applied after replacing the pertinent matrices and vectors with their extended versions. Having established that Algorithm 1 can accommodate multi-lag dependencies, the extension of Algorithm 2 follows, as Algorithm 3 is not affected by the extended state.

Remark 6: One may ponder on the role of graphs in our KRR formulation with the regression coefficient vector obeying a linear dynamical model. Our graph-based formulation is well motivated not only because several physical networks are represented by *graphs* having known connectivity, but also because graphical models are known to represent effectively probabilistic dependencies among nodal vectors. Sure, one could have *a fortiori* assumed knowledge of the needed covariance (or kernel) matrices that are not generally available. Here, we employ kernel matrices that are functions of the graph adjacency matrices. In so doing, we further endow our formulation in the form of regularization terms with graph-related properties that can be present. Those include smoothness, (block) sparsity, low rank, and diffusion effects. Although the emphasis here is on leveraging these properties on graphs the advocated MKriKF approach can be indeed useful in various spatio-temporal estimation tasks involving signals not necessarily evolving over graphs, so long as the underlying (cross-)covariance matrices can become available. Finally, our OKM algorithm leverages the common eigenvectors of the *graph-induced* kernel matrix to reduce complexity. All in all, our graph-based formulation facilitates incorporation of graph-specific prior information.

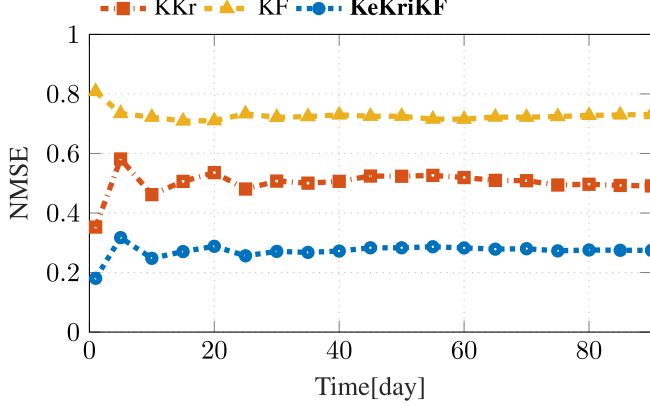
V. SIMULATIONS

This section evaluates the performance of the developed algorithms by means of numerical tests with synthetic and real data. The proposed algorithms are compared with: (i) The least mean-square (LMS) algorithm in [15] with step size μ_{LMS} ; and (ii) the distributed least-squares reconstruction (DLSR) algorithm [30] with step sizes μ_{DLSR} and β_{DLSR} . Both LMS and DLSR can track slowly time-varying B -bandlimited graph signals.

The performance of the aforementioned approaches is quantified through the normalized mean-square error (NMSE)

$$\text{NMSE} := \frac{\mathbb{E} \left[\sum_{\tau=1}^t \|\mathbf{S}_\tau^c (\mathbf{f}_\tau - \hat{\mathbf{f}}_{\tau|\tau})\|_2^2 \right]}{\mathbb{E} \left[\sum_{\tau=1}^t \|\mathbf{S}_\tau^c \mathbf{f}_\tau\|_2^2 \right]}$$

where the expectation is taken over the sample locations, and $\mathbf{S}_\tau^c$ is an $(N - S_\tau) \times N$ matrix comprising the rows of $\mathbf{I}_N$ whose indices are not in $\mathcal{S}_t$, which means that the test set in all experiments is $\mathcal{V} \setminus \mathcal{S}_t, \forall t$. Unless otherwise stated, $\mathcal{S}_t$ is chosen uniformly at random without replacement over $\mathcal{V}$, and kept constant over time; that is, $\mathcal{S}_t = \mathcal{S}$, per t , in order to compare on equal footing the competing algorithms DLSR [30] and LMS [15] that cannot cope with time-varying $\mathcal{S}_t$. For all tests 10 fold cross-validation has been employed. The training set $\mathcal{S}_t$ was split in 10 subsets, out of which 9 were used for training and 1 for validation. The validation error was averaged over the 10 subsets, and the set of parameters exhibiting the smallest error was selected. To reduce the search space, we set $\mu_1 = \mu_2$ and $\rho_\nu = \rho_\chi$. The test set is in all cases $\mathcal{V} \setminus \mathcal{S}_t, \forall t$. Notice that our MKriKF,

Fig. 2. NMSE of function estimates ($\mu_1 = \mu_2 = 1$).

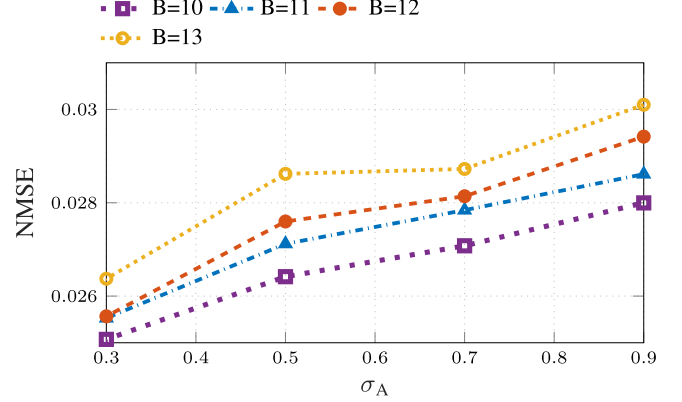
which learns the kernel that “best” fits the data, requires minimal parameter tuning.

A. Numerical Tests on Synthetic Data

To construct a graph, consider the dataset in [4], which contains timestamped messages among students at the University of California, Irvine, exchanged over a social network during 90 days. The sampling interval t is one day. A graph is constructed such that the edge weight $A_{n,n'}(t)$ counts the number of messages exchanged between student n and n' in the k -th month, where $k = 1, 2, 3$ and $30(k-1) + 1 \leq t \leq 30k$. Hence, A_t changes across months. A subset of $N = 310$ users for which A_t corresponds to a connected graph $\forall t$ is selected. At each t , f_t was generated by superimposing a B -bandlimited graph function with $B = 5$ and a spatio-temporally correlated signal. Specifically, $f_t = f_t^{(\nu)} + f_t^{(\chi)} = \sum_{i=1}^5 \gamma_t^i u_t^i + f_t^{(\chi)}$, where $\{\gamma_t^i\}_{i=1}^5 \sim \mathcal{N}(0, 1)$ for all t , while $\{u_t^i\}_{i=1}^5$ denote the eigenvectors associated with the 5 smallest eigenvalues of L_t , and $f_t^{(\chi)}$ is generated according to (9) with $A_{(t,t-1)} = 0.03(A_{t-1} + I_N)$, $\eta \sim \mathcal{N}(0, C_\eta)$, and C_η is a diffusion kernel with $\sigma = 0.5$. Function $f(v_n, t)$ is therefore smooth with respect to the graph and can be interpreted e.g. as the time that the n -th student spends on the specific social network during the t -th day.

The first experiment justifies the proposed decomposition by assessing the impact of dropping either $f_t^{(\nu)}$ or $f_t^{(\chi)}$ from the right hand side of (8). The KriKF algorithm uses diffusion kernels $K_t^{(\nu)}$ and $K_t^{(\chi)}$ with parameters $\sigma = 1.5$ and $\sigma = 0.5$, respectively.

Fig. 2 depicts the NMSE with $S = 217$ for the KeKriKF; the Kalman filter (KF) estimator, which results from setting $f_t^{(\nu)} = 0$ for all t in the KeKriKF; as well as kernel Kriging (KKr), which the KeKriKF reduces to if $f_t^{(\chi)} = 0$ for all t . As observed, KeKriKF, which accounts for both summands in (8), outperforms those algorithms that account for only one of them. Moreover, the low NMSE of KeKriKF in reconstructing the $N - S = 310 - 217 = 93$ unavailable node values reveals that this algorithm is capable of efficiently capturing the spatial as well as the temporal dynamics over time-varying topologies.

Fig. 3. NMSE of KeKriKF for different time-varying graphs ($S = 65$, $\mu_1 = \mu_2 = 1$).

Next, the robustness of KeKriKF is evaluated when the connectivity of $\mathcal{G}_t$, captured by A_t , exhibits abrupt changes over t . Synthetic time-varying networks of size $N = 81$ were generated using the Kronecker product model, which effectively captures properties of real graphs [14]. The prescribed “seed matrix”

$$D_0 := \begin{bmatrix} 1 & 0.1 & 0.7 \\ 0.3 & 0.1 & 0.5 \\ 0 & 1 & 0.1 \end{bmatrix}$$

produces the $N \times N$ matrix $D := D_0 \otimes D_0 \otimes D_0 \otimes D_0$, where $\otimes$ denotes the Kronecker product. An initial adjacency matrix A_0 was constructed with entries $A_{n,n'}(0) \forall n, A_{n,n'}(0) \sim \text{Bernoulli}(D_{n,n'})$ for $n > n'$, and $A_{n,n'}(0) = A_{n',n}(0)$ for $n < n'$. Next, the following time-varying graph model was generated: at each $t_c = 10\kappa$, $\kappa = 1, 2, \dots$, each entry of A_{t_c} changes with probability $p_{n,n'} = \sum_k A_{n,k}(t_c) \sum_l A_{l,n'}(t_c) / \sum_k \sum_l A_{k,l}(t_c)$ as $A_{n,n'}(t_c + 1) = A_{n,n'}(t_c) + |\xi_{n,n'}(t_c)|$ for $n > n'$ where $\xi_{n,n'}(t_c) \sim \mathcal{N}(0, \sigma_A)$ and $A_{n',n}(t_c + 1) = A_{n,n'}(t_c + 1)$ for $n < n'$. This choice of $p_{n,n'}$ is based on the “rich get richer” attribute of real networks, where new connections are formed between nodes with high degree [14]. Moreover, the edge $(v_n, v_{n'})$ is deleted at each $t_d = 20\kappa$, $\kappa = 1, 2, \dots$ with probability 0.1; that is, $A_{n',n}(t_d + 1) = A_{n,n'}(t_d + 1) = 0$, as long as the graph remains connected. By varying σ_A , we obtain different time-varying graphs. A graph function was generated for each time-varying graph as follows

$$f_t = \delta A_t f_{t-1} + \sum_{i=1}^{10} \gamma_t^{(i)} u_t^{(i)} \quad (24)$$

where $\delta = 10^{-2}$ is a forgetting factor, $\sum_{i=1}^{10} \gamma_t^{(i)} u_t^{(i)}$ is a graph-bandlimited component with $\gamma_t^{(i)} \sim \mathcal{N}(0, 1)$, and $\{u_t^{(i)}\}_{i=1}^{10}$ are the eigenvectors associated with the 10 smallest eigenvalues of L_t . Algorithm 1 employs a bandlimited kernel with $\beta = 10^3$ and B for $K_t^{(\nu)}$, a diffusion kernel with $\sigma = 0.5$ for $K_t^{(\chi)}$, and $A_{(t,t-1)} = 10^{-3}(A_{t-1} + I_N)$. Fig. 3 plots the NMSE of the KeKriKF algorithm as a function of σ_A , which determines how

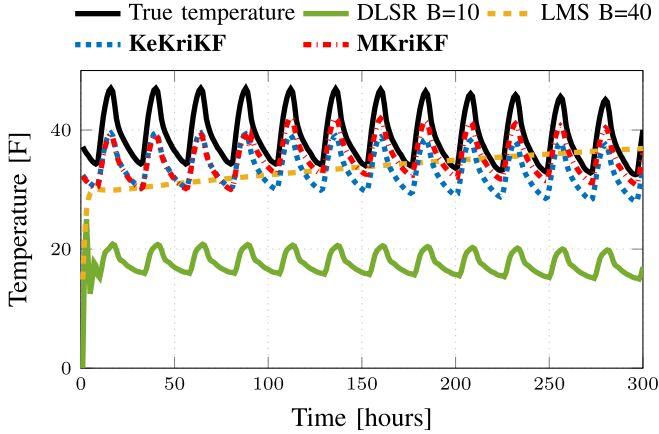


Fig. 4. True and estimated temperature values ($B = 5$, $\mu_{\text{DLSR}} = 1.2$, $\beta_{\text{DLSR}} = 0.5$, $\mu_{\text{LMS}} = 1.5$, $\mu_1 = \mu_2 = 1$, $\rho_\nu = 10^5$, $\rho_\chi = 10^5$).

rapidly the graph changes. As observed, the KeKriKF algorithm can effectively cope with different degrees of time variation.

B. Temperature Prediction

Consider the dataset [1] provided by the National Climatic Data Center, which comprises hourly temperature measurements at $N = 109$ measuring stations across the continental United States in 2010. A time-invariant graph was constructed as in [19], based on geographical distances. The value $f(v_n, t)$ represents the t -th temperature sample recorded at the n -th station. The sampling interval is one hour for the first experiment, and one day for the second.

KeKriKF employs diffusion kernels with parameter $\sigma = 1.8$ for $\mathbf{K}_t^{(\nu)} \mathbf{K}_t^{(\chi)} = 10^{-5} \mathbf{I}_N$, and a transition matrix $\mathbf{A}_{(t,t-1)} = 5 \cdot 10^{-4} (\mathbf{A}_{t-1} + \mathbf{I}_N)$. MKriKF is configured as follows: $\mathcal{D}^{(\nu)}$ contains $M_\nu = 40$ diffusion kernels with parameters $\{\sigma[m]\}_{m=1}^{40}$ with $\sigma[m] \sim \mathcal{N}(2, 0.5)$, $\forall m$; $\mathcal{D}^{(\chi)}$ contains 44 diffusion kernels with parameters $\{\sigma[m]\}_{m=1}^{44}$, where $\sigma[m] \sim \mathcal{N}(1, 0.2)$, $\forall m$, and an identity kernel $\mathbf{K}^{(\chi)}[45] = \mathbf{I}_N$.

Fig. 4 depicts the true temperature along with its estimates for a station n that is not sampled, meaning $n \notin \mathcal{S}$, with $S = 44$. Clearly, KeKriKF accurately tracks the temperature by exploiting spatial and temporal dynamics, but MKriKF outperforms KeKriKF by learning those dynamics from the data. The random sampling set selection heavily affects performance of the LMS algorithm; for adaptive selection of $\mathcal{S}$ see [15].

Fig. 5 compares the NMSE of all considered approaches for $S = 44$. Observe the superior performance of the proposed reconstruction methods, which in this scenario exhibit roughly the same NMSE.

C. GDP Prediction

The next dataset is provided by the World Bank Group [2], and comprises gross domestic product (GDP) per capita for $N = 127$ countries for the years 1960–2016. A time-invariant graph was constructed using the correlation between the GDP of different countries for the first 25 years. The graph function

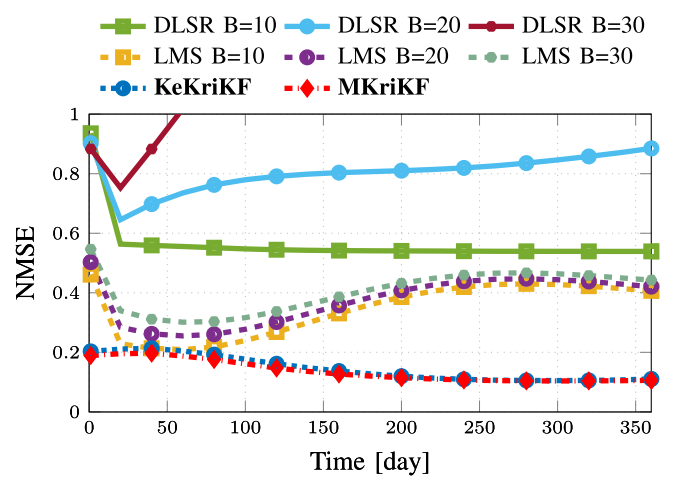


Fig. 5. NMSE of temperature estimates ($\mu_{\text{DLSR}} = 1.6$, $\beta_{\text{DLSR}} = 0.5$, $\mu_{\text{LMS}} = 1.5$, $\rho_\nu = 10^5$, $\rho_\chi = 10^5$).

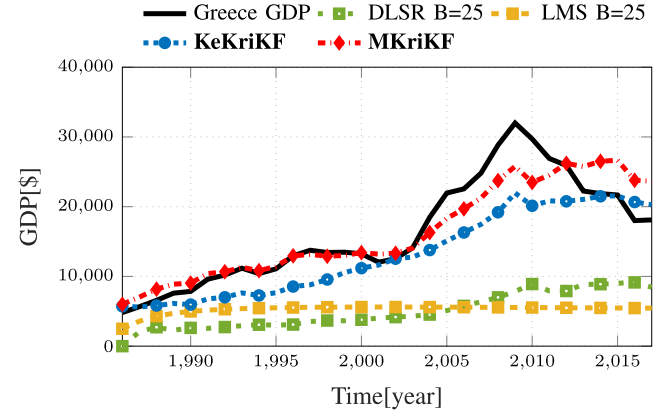


Fig. 6. Greece GDP values along with the estimated ones ($S = 38$, $\mu_{\text{DLSR}} = 1.6$, $\beta_{\text{DLSR}} = 0.4$, $\mu_{\text{LMS}} = 1.2$, $\rho_\nu = 10^4$, $\rho_\chi = 10^4$).

$f(v_n, t)$ denotes the GDP reported at the n -th country and t -th year for $t = 1985, \dots, 2016$.

The graph Fourier transform of the GDP in the first 25 years defined as $\tilde{f}_n := \mathbf{u}_n^\top \mathbf{f} \forall n$, where $\mathbf{u}_n$ denotes the n -th eigenvector of the Laplacian matrix; see [26], shows that the graph frequencies $\tilde{f}_k$ take small values for $4 < k < 123$, and large values otherwise. Motivated by the aforementioned observation, the KeKriKF is configured with a band-reject kernel $\mathbf{K}^{(\nu)}$ with $k = 6$, $l = 6$, $\beta = 15$; see Table I, $\mathbf{K}^{(\chi)} = 10^{-3} \mathbf{I}_N$, and $\mathbf{A}_{(t,t-1)} = 10^{-5} (\mathbf{A}_{t-1} + \mathbf{I}_N)$. MKriKF adopts a $\mathcal{D}^{(\nu)}$ with $M_\nu = 16$ band-reject kernels with $k \in [2, 5]$, $l \in [1, 4]$, $\beta = 15$, and a $\mathcal{D}^{(\chi)}$ with 60 diffusion kernels with parameters $\{\sigma[m]\}_{m=1}^{60}$, where $\sigma[m] \sim \mathcal{N}(2, 0.5)$, $\forall m$, and an identity kernel $\mathbf{K}^{(\chi)}[61] = \mathbf{I}_N$.

Fig. 6 depicts the actual GDP as well as its estimates for Greece, which is not contained in the sampled countries. Clearly, both MKriKF and KeKriKF, track the GDP evolution over the years with greater accuracy than the considered alternatives. This is expected because the graph function does not adhere to the graph bandlimited model assumed by DLSR and LMS.

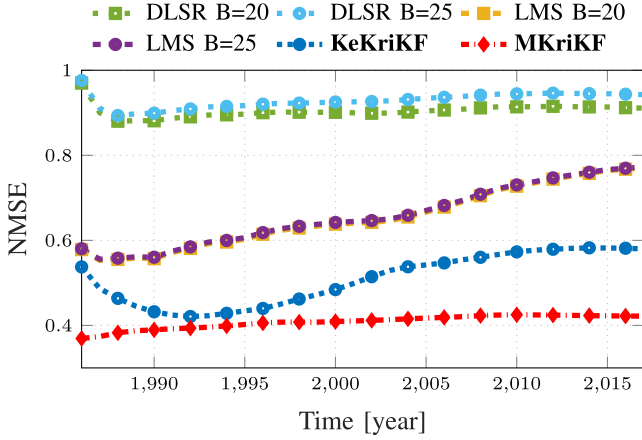


Fig. 7. NMSE of GDP estimates ($S = 38$, $\mu_{\text{DLSR}} = 1.6$, $\beta_{\text{DLSR}} = 0.4$, $\mu_{\text{LMS}} = 1.6$, $\rho_\nu = 10^4$, $\rho_\chi = 10^4$).

Fig. 7 reports NMSE over time, where the proposed algorithms achieve the smallest NMSE. The data-driven MKriKF outperforms KeKriKF, which is configured manually.

D. Network Delay Prediction

The last dataset records measurements of path delays on the Internet2 backbone [3]. The network comprises 9 end-nodes and 26 directed links. The delays are available for $N = 70$ paths at every minute. The paths connect origin-destination nodes by a series of links described by the path-link routing matrix $\Pi \in \{0, 1\}^{N \times 26}$, whose (n, l) entry is $\Pi_{n,l} = 1$ if path n' traverses link l , and 0 otherwise. A graph is constructed with each vertex corresponding to one of these paths, and with the time-invariant adjacency matrix $A \in \mathbb{R}^{N \times N}$ given by

$$A_{n,n'} = \frac{\sum_{l=1}^{26} \Pi_{n,l} \Pi_{n',l}}{\sum_{l=1}^{26} \Pi_{n,l} + \sum_{l=1}^{26} \Pi_{n',l} - \sum_{l=1}^{26} \Pi_{n,l} \Pi_{n',l}} \quad (25)$$

for $n, n' = 1, \dots, N$, $n \neq n'$. Expression (25) was selected to assign a greater weight to edges connecting vertices whose associated paths share a large number of links. This is intuitively reasonable since paths with common links usually experience similar delays [7]. Function $f(v_n, t)$ denotes the delay in milliseconds measured at the n -th path and t -th minute.

The KeKriKF algorithm employs a diffusion kernel with parameter $\sigma = 2.5$ for $K_t^{(\nu)}$, $K_t^{(\chi)} = 0.002I_N$, and $A_{(t,t-1)} = 0.005(A_{t-1} + I_N)$. The MKriKF is configured as follows: $\mathcal{D}^{(\nu)}$ contains $M_\nu = 40$ diffusion kernels with parameters $\{\sigma[m]\}_{m=1}^{40}$ with $\sigma[m] \sim \mathcal{N}(4, 0.5)$, $\forall m$; $\mathcal{D}^{(\chi)}$ contains $M_\chi = 60$ diffusion kernels with parameters $\{\sigma[m]\}_{m=1}^{60}$ with $\sigma[m] \sim \mathcal{N}(1, 0.1)$, $\forall m$, and an identity kernel $K^{(\chi)}[61] = I_N$.

Fig. 8 depicts the NMSE when $S = 20$. KeKriKF and MKriKF are seen to outperform competing methods. Using the parameters of Fig. 8, the MKriKF algorithm is tested for S_t chosen at random per slot t with $S_t = S$, $\forall t$, but $S_t \neq S_{t'}, \forall t \neq t'$. Fig. 9 shows the NMSE of MKriKF with variable S and as expected the performance improves as the number of samples increases. For the same configuration, Fig. 10 depicts the NMSE

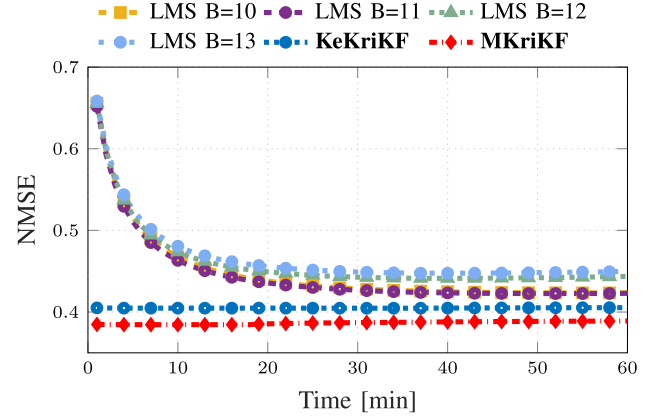


Fig. 8. NMSE of network delay estimates ($\mu_{\text{LMS}} = 1.5$, $c = 0.0005$, $\rho_\nu = 100$, $\rho_\chi = 100$, $\mu_1 = \mu_2 = 1$).

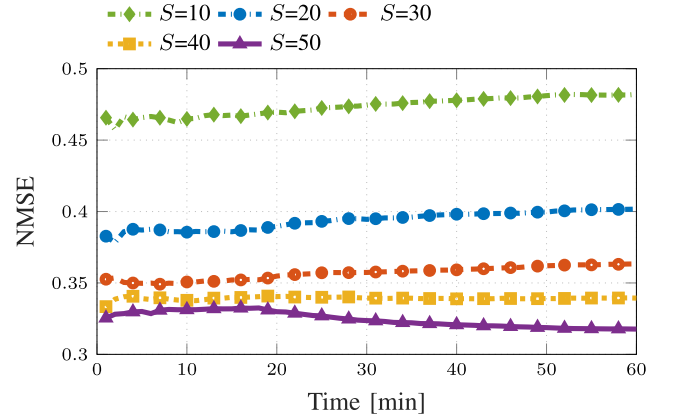


Fig. 9. NMSE of MKriKF with varying sampling set size ($\mu_1 = \mu_2 = 1$, $\rho_\nu = 100$, $\rho_\chi = 100$).

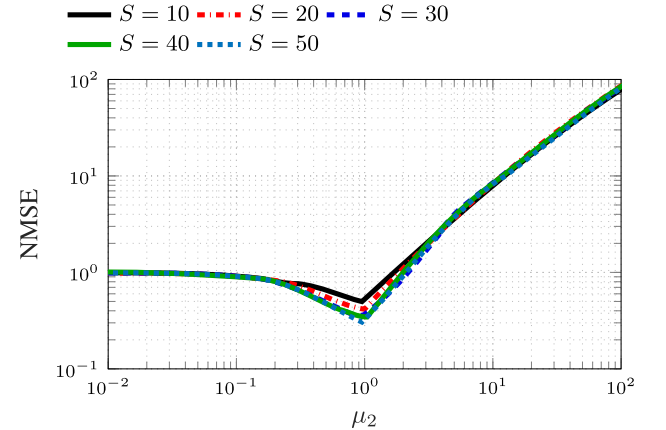


Fig. 10. NMSE of MKriKF for different μ_2 ($\mu_1 = 1$, $\rho_\nu = 100$, $\rho_\chi = 100$).

as μ_2 varies, and shows that the minimum NMSE for all S is at $\mu_2 = 1$.

Finally, the proposed MKriKF will be evaluated in tracking the delay over the network from $S = 56$ randomly sampled path delays. To that end, delay maps are traditionally employed, which depict the network delay per path over time and enable operators to perform troubleshooting; see also [18]. The paths

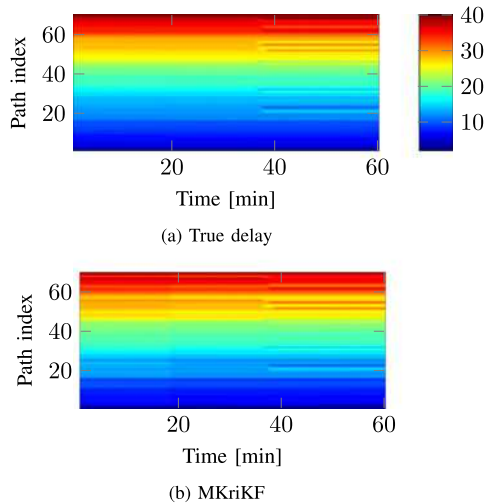


Fig. 11. True and estimated network delay map for $N = 70$ paths ($\rho_\chi = 100$, $\mu_1 = \mu_2 = 1$). (a) True delay (b) MKriKF.

for the delay maps in Fig. 11 are sorted in increasing order of the true delay at $t = 1$. Clearly, the delay map recovered by MKriKF in Fig. 11 (b) visually resembles the true delay map in Fig. 11 (a).

VI. CONCLUSION

This paper introduced online estimators to reconstruct dynamic functions over (possibly dynamic) graphs. In this context, the function to be estimated was decomposed in two parts: one capturing the spatial dynamics, and the other jointly modeling spatio-temporal dynamics by means of a state-space model. A novel kernel kriged Kalman filter was developed using a deterministic RKHS approach. To accommodate scenarios with limited prior information, an online multi-kernel learning technique was also developed to allow tracking of the spatio-temporal dynamics of the graph function. The structure of Laplacian kernels was exploited to achieve low computational complexity. Through numerical tests with synthetic as well as real-data, the novel algorithms were observed to perform markedly better than existing alternatives. Future work includes distributed implementations of the proposed algorithms, data-driven learning of $\mathbf{A}_{(t,t-1)}$, and exploring nonlinear dynamical models such as the sampled Brownian motion, the extended KF, unscented KF, or particle filters.

REFERENCES

- [1] "1981–2010 U.S. climate normals," [Online]. Available: <https://www.ncdc.noaa.gov/data-access/land-based-station-data/land-based-datasets/climate-normals/1981-2010-normals-data>, Accessed on: Sep. 2016.
- [2] "GDP per capita (current US)," [Online]. Available: <https://data.worldbank.org/indicator/NY.GDP.PCAP.CD>, Accessed on: Sep. 2017.
- [3] "One-way ping internet2," [Online]. Available: <http://software.internet2.edu/owamp/>, Accessed on: Sep. 2017.
- [4] "Snap temporal networks: Collegemsg," [Online]. Available: <http://snap.stanford.edu/data/CollegeMsg.html>, Accessed Sep. 2017.
- [5] A. Anis, A. Gadde, and A. Ortega, "Efficient sampling set selection for bandlimited graph signals using graph spectral proxies," *IEEE Trans. Signal Process.*, vol. 64, no. 14, pp. 3775–3789, Jul. 2016.
- [6] D. Bertsekas, *Nonlinear Programming*. Belmont, MA, USA: Athena Scientific, 1999.
- [7] D. B. Chua, E. D. Kolaczyk, and M. Crovella, "Network kriging," *IEEE J. Sel. Areas Commun.*, vol. 24, no. 12, pp. 2263–2272, Dec. 2006.
- [8] I. Csizsár and G. Tusnádý, "Information geometry and alternating minimization procedures," *Stat. Decis., Supplement Issue, 1*, pp. 205–237, 1984.
- [9] P. A. Forero, K. Rajawat, and G. B. Giannakis, "Prediction of partially observed dynamical processes over networks via dictionary learning," *IEEE Trans. Signal Process.*, vol. 62, no. 13, pp. 3305–3320, Jul. 2014.
- [10] V. N. Ioannidis, A. N. Nikolakopoulos, and G. B. Giannakis, "Semi-parametric graph kernel-based reconstruction," in *Proc. Global Conf. Signal Inf. Process.*, Montreal, QC, Canada, Nov. 2017, pp. 588–592.
- [11] V. N. Ioannidis, D. Romero, and G. B. Giannakis, "Inference of spatiotemporal processes over graphs via kernel kriged Kalman filtering," in *Proc. Eur. Signal Process. Conf.*, Kos, Greece, Aug. 2017, pp. 1679–1683.
- [12] E. D. Kolaczyk, *Statistical Analysis of Network Data: Methods and Models*. New York, NY, USA: Springer, 2009.
- [13] R. I. Kondor and J. Lafferty, "Diffusion kernels on graphs and other discrete structures," in *Proc. Int. Conf. Mach. Learn.*, Sydney, NSW, Australia, Jul. 2002, pp. 315–322.
- [14] J. Leskovec, D. Chakrabarti, J. Kleinberg, C. Faloutsos, and Z. Ghahramani, "Kronecker graphs: An approach to modeling networks," *J. Mach. Learn. Res.*, vol. 11, pp. 985–1042, Feb. 2010.
- [15] P. D. Lorenzo, S. Barbarossa, P. Banelli, and S. Sardellitti, "Adaptive least mean-square estimation of graph signals," *IEEE Trans. Signal Inf. Process. Netw.*, vol. 2, no. 4, pp. 555–568, Sep. 2016.
- [16] H. Lütkepohl, *New Introduction to Multiple Time Series Analysis*. New York, NY, USA: Springer, 2005.
- [17] K. V. Mardia, C. Goodall, E. J. Redfern, and F. J. Alonso, "The kriged Kalman filter," *Test*, vol. 7, no. 2, pp. 217–282, 1998.
- [18] K. Rajawat, E. Dall'Anese, and G. B. Giannakis, "Dynamic network delay cartography," *IEEE Trans. Inf. Theory*, vol. 60, no. 5, pp. 2910–2920, Mar. 2014.
- [19] D. Romero, V. N. Ioannidis, and G. B. Giannakis, "Kernel-based reconstruction of space-time functions on dynamic graphs," *IEEE J. Sel. Topics Signal Process.*, vol. 11, no. 6, pp. 1–14, Sep. 2017.
- [20] D. Romero and G. Leus, "Wideband spectrum sensing from compressed measurements using spectral prior information," *IEEE Trans. Signal Process.*, vol. 61, no. 24, pp. 6232–6246, Dec. 2013.
- [21] D. Romero, M. Ma, and G. B. Giannakis, "Kernel-based reconstruction of graph signals," *IEEE Trans. Signal Process.*, vol. 65, no. 3, pp. 764–778, Feb. 2017.
- [22] I. D. Schizas, G. B. Giannakis, S. I. Roumeliotis, and A. Ribeiro, "Consensus in ad hoc WSNs with noisy links—Part II: Distributed estimation and smoothing of random signals," *IEEE Trans. Signal Process.*, vol. 56, no. 4, pp. 1650–1666, Apr. 2008.
- [23] B. Schölkopf and A. J. Smola, *Learning with Kernels: Support Vector Machines, Regularization, Optimization, and Beyond*. Cambridge, MA, USA: MIT Press, 2002.
- [24] S. Segarra, A. G. Marques, G. Leus, and A. Ribeiro, "Reconstruction of graph signals through percolation from seeding nodes," *IEEE Trans. Signal Process.*, vol. 64, no. 16, pp. 4363–4378, Aug. 2016.
- [25] Y. Shen, B. Baingana, and G. B. Giannakis, "Nonlinear structural vector autoregressive models for inferring effective brain network connectivity," arXiv:1610.06551v1, 2016.
- [26] D. I. Shuman, S. K. Narang, P. Frossard, A. Ortega, and P. Vandergheynst, "The emerging field of signal processing on graphs: Extending high-dimensional data analysis to networks and other irregular domains," *IEEE Signal Process. Mag.*, vol. 30, no. 3, pp. 83–98, May 2013.
- [27] A. J. Smola and R. I. Kondor, "Kernels and regularization on graphs," in *Learning Theory and Kernel Machines*, New York, NY, USA: Springer, 2003, pp. 144–158.
- [28] G. Strang and K. Borre, *Linear Algebra, Geodesy, and GPS*. Philadelphia, PA, USA: SIAM, 1997.
- [29] D. Thanou, D. I. Shuman, and P. Frossard, "Learning parametric dictionaries for signals on graphs," *IEEE Trans. Signal Process.*, vol. 62, no. 15, pp. 3849–3862, Aug. 2014.
- [30] X. Wang, M. Wang, and Y. Gu, "A distributed tracking algorithm for reconstruction of graph signals," *IEEE J. Sel. Topics Signal Process.*, vol. 9, no. 4, pp. 728–740, Feb. 2015.
- [31] K. K. Wike and N. Cressie, "A dimension-reduced approach to space-time Kalman filtering," *Biometrika*, vol. 86, pp. 815–829, 1999.
- [32] D. Zhou and B. Schölkopf, "A regularization framework for learning from graph data," in *Proc. ICML Workshop Stat. Relational Learn. Connections Other Fields*, Banff, AB, Canada, Jul. 2004, vol. 15, pp. 67–68.



analytics, and network science. He was the recipient of the Performance Excellence award.



interests include signal processing, communications, and machine learning.

Vassilis N. Ioannidis (S'16) received the Diploma in electrical and computer engineering from the National Technical University of Athens, Athens, Greece, in 2015, and the M.Sc. degree in electrical engineering from the University of Minnesota, Minneapolis, MN, USA, in 2017. He is currently working toward the Ph.D. degree at the Department of Electrical and Computer Engineering, University of Minnesota. From 2014 to 2015, he worked as a Middleware Consultant for Oracle in Athens, Greece. His research interests include machine learning, big data



Georgios B. Giannakis (F'97) received the Diploma in electrical engineering from the National Technical University of Athens, Athens, Greece, 1981. From 1982 to 1986, he was with the University of Southern California, where he received the MSc. degree in electrical engineering, 1983, the MSc. degree in mathematics, 1986, and the Ph.D. degree in electrical engineering, 1986.

He was with the University of Virginia from 1987 to 1998, and since 1999, he has been a Professor with the University of Minnesota, Minneapolis, MN, USA, 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 interests include communications, networking and statistical signal processing—subjects on which he has authored or coauthored more than 400 journal papers, 700 conference papers, 25 book chapters, 2 edited books and 2 research monographs (h-index 125). His current research interest focuses on learning from big data, wireless cognitive radios, and network science with applications to social, brain, and power networks with renewables. 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. He is the (co-) inventor of 30 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 was also the recipient of the 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).