

OPTIMAL REDUCED MODEL ALGORITHMS FOR DATA-BASED STATE ESTIMATION*

ALBERT COHEN[†], WOLFGANG DAHMEN[‡], RONALD DEVORE[§], JALAL FADILI[¶],
OLGA MULA^{||}, AND JAMES NICHOLS[#]

Abstract. Reduced model spaces, such as reduced bases and polynomial chaos, are linear spaces V_n of finite dimension n which are designed for the efficient approximation of certain families of parametrized PDEs in a Hilbert space V . The manifold $\mathcal{M}$ that gathers the solutions of the PDE for all admissible parameter values is globally approximated by the space V_n with some controlled accuracy ε_n , which is typically much smaller than when using standard approximation spaces of the same dimension such as finite elements. Reduced model spaces have also been proposed in [Y. Maday et al., *Internat. J. Numer. Methods Engrg.*, 102 (2015), pp. 933–965] as a vehicle to design a simple linear recovery algorithm of the state $u \in \mathcal{M}$ corresponding to a particular solution instance when the values of parameters are unknown but a set of data is given by m linear measurements of the state. The measurements are of the form $\ell_j(u)$, $j = 1, \dots, m$, where the ℓ_j are linear functionals on V . The analysis of this approach in [P. Binev et al., *SIAM/ASA J. Uncertain. Quantif.*, 5 (2017), pp. 1–29] shows that the recovery error is bounded by $\mu_n \varepsilon_n$, where $\mu_n = \mu(V_n, W)$ is the inverse of an inf-sup constant that describe the angle between V_n and the space W spanned by the Riesz representers of $(\ell_1, \dots, \ell_m)$. A reduced model space which is efficient for approximation might thus be ineffective for recovery if μ_n is large or infinite. In this paper, we discuss the existence and effective construction of an optimal reduced model space for this recovery method. We extend our search to affine spaces which are better adapted than linear spaces for various purposes. Our basic observation is that this problem is equivalent to the search of an optimal affine algorithm for the recovery of $\mathcal{M}$ in the worst case error sense. This allows us to perform our search by a convex optimization procedure. Numerical tests illustrate that the reduced model spaces constructed from our approach perform better than the classical reduced basis spaces.

Key words. reduced models, optimal recovery, sensing, parametrized PDEs, convex optimization

AMS subject classifications. 65M32, 65M15, 68U20, 65D99

DOI. 10.1137/19M1255185

1. Introduction.

1.1. Background and context. State estimation refers to the general problem of approximately recovering the true state of a physical system of interest from incomplete data. This task is ubiquitous in applied sciences and engineering. One can draw a distinction between two different application scenarios:

*Received by the editors April 8, 2019; accepted for publication (in revised form) July 13, 2020; published electronically November 24, 2020.

<https://doi.org/10.1137/19M1255185>

Funding: The work of the third author was supported by the Institut Universitaire de France, the ERC Adv grant BREAD, the EMERGENCE grant project “Models and Measures” from the Paris City Council, the National Science Foundation grants DMS 15-21067, DMS 18-17603, DMS 172097 and the Office of Naval Research grants N00014-17-1-2908, N00014-16-1-2706. The work of the first, second, third, fifth and sixth authors were also supported as visitors by the Isaac Newton Institute at Cambridge University.

[†]Sorbonne Université, Paris, 75005, France (albert.cohen@sorbonne-universite.fr).

[‡]Department of Mathematics, University of South Carolina, Columbia, SC 29208 USA (wolfgang.anton.dahmen@gmail.com).

[§]Texas A&M University, College Station, TX 77843 USA (rdevore@math.tamu.edu).

[¶]Normandie Université, ENSICAEN, CNRS, Caen, 14050, France (Jalal.Fadili@ensicaen.fr).

^{||}CEREMADE, CNRS, UMR 7534, Université Paris-Dauphine, PLS University, 75016, Paris, France; Inria, Sorbonne Université and CNRS, Paris, France (mula@ceremade.dauphine.fr).

[#]Australian National University, ACT 0200, Australia (james.nichols@anu.edu.au).

- (i) The physical properties of the states, sometimes referred to as background information, are approximately modeled by a nonlinear dynamical system which by itself is neither sufficiently accurate nor stable to warrant reliable predictions. This is typically the case in weather prediction, climatology, or generally in atmospheric research. One therefore utilizes observational data to *correct* the model-based predictions, ideally in real time. Such correction mechanisms are often based on statistical hypotheses such as Gaussianity of error distributions. One important approach is *ensemble Kalman filtering* which can be viewed as a recursive Bayesian estimation based on Monte Carlo approximations to the first and second moments of the error distributions [16, 18]. A second class of methods are so-called variational data assimilation schemes like 3 dimensional (3D)-VAR or 4 dimensional (4D)-VAR [19]. The state predicted by the model is then corrected by minimizing a quadratic cost functional involving inverse covariance matrices for the background model error and observation error. In this first scenario, the error bounds between the exact and estimated state are typically expressed in an average sense, based on the accepted simplified statistical model assumptions.
- (ii) The physical states of interest are reliably described in terms of a *parameter dependent* family of PDEs which for each parameter instance can be computed within a desired target accuracy. The states are therefore elements of the associated *solution manifold* that consists of all solutions to the PDE as parameters vary. The task is then to estimate a state in (or near to) the solution manifold from only a finite number of measurements generated through a *fixed* number of sensors. A classical example is to estimate a pressure field of a porous media flow from a finite number of pressure head measurements. The parametric model then could arise from a Karhunen–Loève expansion of a random field of permeability coefficients in Darcy’s law, and may thus involve a large or even infinite number of parameters. Problems of this type have been investigated over the past decade in the context of uncertainty quantification. Again, a prominent approach is Bayesian inversion where prior information is given in terms of a probability distribution for the parameter, inducing a probability distribution for the state [25]. The objective is then to approximate the posterior probability distribution of states given the data. High dimensionality renders such methods computationally expensive. Alternatively, state estimation can be formulated as a constrained optimization problem. For instance, one could minimize the deviation of state measurements over the solution manifold asking for probabilistic or deterministic error bounds. In practice, one typically chooses first a sufficiently fine discretization of the *high fidelity* continuum model which then gives rise to a large scale (discrete) constrained nonconvex optimization problem that needs to be solved for each instance of data. Ill-posedness of the inversion task necessitates adding regularization terms which introduce a further ambiguous bias. Reduced models are used to alleviate the possibly prohibitive cost of the numerous forward simulations that are needed in the descent method. A central issue is then to judiciously switch between the high fidelity model, given in terms of the fine scale discretization, and the low fidelity reduced model; see [27].

In this article, we consider scenario (ii) but pursue a different approach taking up on recent work in [3, 20]. Although it can be formulated without any reference to a statistical model, it has conceptual similarities with the 3D- and 4D-Var variational approach invoked for scenario (i); see [17] and [26] for such connections. In contrast

to Bayesian inversion, this approach yields deterministic error bounds expressed in a worst case sense over the solution manifold, which is the primary interest in this paper. Specifically, we follow [3] and formulate state estimation as an *optimal recovery problem*; see, e.g., [21]. This allows us to formulate optimality benchmarks that steer our development of recovery algorithms.

1.2. Mathematical formulation of the state estimation problem. The *sensing* or *recovery* problem studied in this paper is formulated in a Hilbert space V equipped with some norm $\|\cdot\|$ and inner product $\langle\cdot,\cdot\rangle$: we want to recover an approximation to an unknown function $u \in V$ from data given by m linear measurements

$$(1.1) \quad \ell_i(u), \quad i = 1, \dots, m,$$

where the ℓ_i are m linearly independent bounded linear functionals over V . This problem appears in many different settings. The particular one that motivates our work is the case where $u = u(y)$ represents the *state* of a physical system described as a solution to a parametric PDE

$$(1.2) \quad \mathcal{P}(u, y) = 0$$

for some unknown finite or infinite dimensional parameter vector $y = (y_j)_{j \geq 1}$ picked from some admissible set Y . The ℓ_i are a mathematical model for sensors that capture some partial information on the unknown solution $u(y) \in V$.

Denoting by $\omega_i \in V$ the Riesz representers of the ℓ_i , such that $\ell_i(v) = \langle \omega_i, v \rangle$ for all $v \in V$, and defining

$$(1.3) \quad W := \text{span}\{\omega_1, \dots, \omega_m\},$$

the measurements are equivalently represented by

$$(1.4) \quad w = P_W u,$$

where P_W is the orthogonal projection from V onto W . A *recovery algorithm* is a computable map

$$(1.5) \quad A : W \rightarrow V$$

and the approximation to u obtained by this algorithm is

$$(1.6) \quad u^* = A(w) = A(P_W u).$$

The construction of A should be based on the available prior information that describes the properties of the unknown u , and the evaluation of its performance needs to be defined in some precise sense. Two distinct approaches are usually followed:

- In the *deterministic setting*, the sole prior information is that u belongs to the set

$$(1.7) \quad \mathcal{M} := \{u(y) : y \in Y\}$$

of all possible solutions. The set $\mathcal{M}$ is sometimes called the *solution manifold*. The performance of an algorithm A over the class $\mathcal{M}$ is usually measured by the “worst case” reconstruction error

$$(1.8) \quad E_{\text{wc}}(A, \mathcal{M}) = \sup\{\|u - A(P_W u)\| : u \in \mathcal{M}\}.$$

The problem of finding an algorithm that minimizes $E_{\text{wc}}(A)$ is called *optimal recovery*. It has been extensively studied for convex sets $\mathcal{M}$ that are balls of smoothness classes [5, 21, 22], which is not the case for (1.7).

- In the *stochastic setting*, the prior information on u is described by a probability distribution p on V , which is supported on $\mathcal{M}$, typically induced by a probability distribution on Y that is assumed to be known. It is then natural to measure the performance of an algorithm in an averaged sense, for example, through the mean square error

$$(1.9) \quad E_{\text{ms}}(A, p) = \mathbb{E}(\|u - A(P_W u)\|^2) = \int_V \|u - A(P_W u)\|^2 dp(u).$$

This stochastic setting is the starting point for *Bayesian estimation* methods [13]. Let us observe that for any algorithm A one has $E_{\text{ms}}(A, p) \leq E_{\text{wc}}(A, \mathcal{M})^2$.

1.3. Optimal algorithms. The present paper concentrates on the deterministic setting according to the above distinction, although some remarks will be given on the analogies with the stochastic setting. In this setting, the benchmark for the performance of recovery algorithms is given by

$$E_{\text{wc}}^*(\mathcal{M}) = \inf_A E_{\text{wc}}(A, \mathcal{M}),$$

where the infimum is taken over all possible maps A .

There is a simple mathematical description of an optimal map that meets this benchmark. For any bounded set $S \subset V$ we define its *Chebyshev ball* as the smallest closed ball that contains S . The *Chebyshev ball radius and center* denoted by $\text{rad}(S)$ and $\text{cen}(S)$ are the radius and center of this ball. Since the information that we have on u is that it belongs to the set

$$(1.10) \quad \mathcal{M}_w := \mathcal{M} \cap V_w, \quad V_w := \{v \in V : P_W v = w\} = w + W^\perp,$$

where $W^\perp$ is the orthogonal complement of W in V , it follows that an optimal reconstruction map A_{wc}^* for the worst case error is given by

$$(1.11) \quad A_{\text{wc}}^*(w) = \text{cen}(\mathcal{M}_w),$$

because the Chebyshev ball center of $\mathcal{M}_w$ minimizes the quantity $\sup\{\|u - v\| : u \in \mathcal{M}_w\}$ among all $v \in V$. The worst case error is therefore given by

$$(1.12) \quad E_{\text{wc}}^*(\mathcal{M}) = E_{\text{wc}}(A_{\text{wc}}^*, \mathcal{M}) = \sup\{\text{rad}(\mathcal{M}_w) : w \in P_W(\mathcal{M})\}.$$

Note that the map A_{wc}^* is also optimal among all algorithms for each $\mathcal{M}_w$, $w \in P_W(\mathcal{M})$, since

$$(1.13) \quad E_{\text{wc}}(A_{\text{wc}}^*, \mathcal{M}_w) = \min_A E_{\text{wc}}(A, \mathcal{M}_w) = \text{rad}(\mathcal{M}_w), \quad w \in P_W(\mathcal{M}).$$

However, there may exist other maps A such that $E_{\text{wc}}(A, \mathcal{M}) = E_{\text{wc}}^*(\mathcal{M})$, since we also supremize over $w \in P_W(\mathcal{M})$.

1.4. Linear and affine algorithms based on reduced models. In practice the above map A_{wc}^* cannot be easily constructed due to the fact that the solution manifold $\mathcal{M}$ is a high dimensional and geometrically complex object. One is therefore interested in designing “suboptimal yet good” recovery algorithms and analyzing their performance.

One vehicle for constructing linear recovery mappings A is to use *reduced modeling*. Generally speaking, reduced models consist of linear spaces $(V_n)_{n \geq 0}$ with increasing dimension $\dim(V_n) = n$ which uniformly approximate the solution manifold in the sense that

$$(1.14) \quad \text{dist}(\mathcal{M}, V_n) := \max_{u \in \mathcal{M}} \min_{v \in V_n} \|u - v\| \leq \varepsilon_n,$$

where

$$(1.15) \quad \varepsilon_0 \geq \varepsilon_1 \geq \cdots \geq \varepsilon_n \geq \cdots \geq 0$$

are known tolerances. Instances of reduced models for parametrized families of PDEs with provable accuracy are provided by polynomial approximations in the y variable [10, 11] or reduced bases [7, 24, 23]. The construction of a reduced model is typically done offline, possibly using a large training set of instances of $u \in \mathcal{M}$ called *snapshots*. The offline stage potentially has a high computational cost. Once this is done, the online cost of recovering $u^* = A(w)$ from any data w using this reduced model should, in contrast, be moderate.

In [20], a simple reduced-model-based recovery algorithm was proposed in terms of the map

$$(1.16) \quad A_n(w) := \operatorname{argmin}\{\text{dist}(v, V_n) : v \in V_w\},$$

which is well-defined provided that $V_n \cap W^\perp = \{0\}$. It turns out that A_n is a linear mapping and so these algorithms are linear. This approach is called the parametrized-background data-weak method, however, we follow the terminology introduced in [3], referring to an algorithm of the form A_n as a *one-space algorithm*. In the latter, it was shown that A_n has a simple interpretation in terms of the cylinder

$$(1.17) \quad \mathcal{K}_n := \{v \in V : \text{dist}(v, V_n) \leq \varepsilon_n\}$$

that contains the solution manifold $\mathcal{M}$. Namely, the algorithm A_n is also given by

$$(1.18) \quad A_n(w) = \text{cen}(\mathcal{K}_n, w), \quad \mathcal{K}_{n,w} := \mathcal{K}_n \cap V_w,$$

and the map is shown to be optimal when $\mathcal{M}$ is replaced by the simpler containment set $\mathcal{K}_n$, that is,

$$A_n = \operatorname{argmin}_{A: W \rightarrow V} E_{\text{wc}}(A, \mathcal{K}_n).$$

The substantial advantage of this approach is that, in contrast to A_{wc}^* , the map A_n can be easily computed by solving simple least-squares minimization problems which amount to finite linear systems. In turn A_n is a linear map from W to V . This map depends on V_n and W , but not on ε_n in view of (1.16). We refer to A_n as the one-space algorithm based on the space V_n .

This algorithm satisfies the performance bound

$$(1.19) \quad \|u - A_n(P_W u)\| \leq \mu_n \text{dist}(u, V_n \oplus (V_n^\perp \cap W)) \leq \mu_n \text{dist}(u, V_n) \leq \mu_n \varepsilon_n,$$

where the last inequality holds when $u \in \mathcal{M}$. Here

$$(1.20) \quad \mu_n = \mu(V_n, W) := \max_{v \in V_n} \frac{\|v\|}{\|P_W v\|}$$

is the inverse of the inf-sup constant $\beta_n := \min_{v \in V_n} \max_{w \in W} \frac{\langle v, w \rangle}{\|v\| \|w\|}$ which describes the angle between V_n and W . In particular $\mu_n = \infty$ in the event where $V_n \cap W^\perp$ is nontrivial.

An important observation is that the one-space algorithm (1.16) has a simple extension to the setting where V_n is an affine space rather than a linear space, namely, when

$$(1.21) \quad V_n = \bar{u} + \tilde{V}_n$$

with $\tilde{V}_n$ a linear space of dimension n and $\bar{u}$ a given offset that is known to us.

At a first sight, affine spaces do not bring any significant improvement in terms of approximating the solution manifold, due to the following elementary observation: if $\mathcal{M}$ is approximated with accuracy ε by an n -dimensional affine space V_n given by (1.21), it is also approximated with accuracy $\tilde{\varepsilon} \leq \varepsilon$ by the $(n+1)$ -dimensional linear space

$$(1.22) \quad \tilde{V}_{n+1} := V_n \oplus \mathbb{R}\bar{u}.$$

However, the choice of an affine reduced model may significantly improve the performance of the one-space algorithm in the case where the parametric solution $u(y)$ is a “small perturbation” of a nominal solution $\bar{u} = u(\bar{y})$ for some $\bar{y} \in Y$ in the sense that

$$(1.23) \quad \text{diam}(\mathcal{M}) \ll \|\bar{u}\|.$$

Indeed, suppose in addition that $\bar{u}$ is badly aligned with respect to the measurement space W in the sense that

$$(1.24) \quad \|P_W \bar{u}\| \ll \|\bar{u}\|.$$

In such a case, any linear space V_n that is well tailored to approximating the solution manifold (for example, a reduced basis space) will contain a direction close to that of $\bar{u}$ and thus, we will have that $\mu_n \gg 1$, rendering the reconstruction by the linear one-space method much less accurate than the approximation error by V_n . The use of the affine mapping (1.21) has the advantage of eliminating the bad direction $\bar{u}$ since μ_n will now be computed with respect to the linear part $\tilde{V}_n$.

A further perspective, currently under investigation, is to agglomerate local affine models in order to generate a *nonlinear* reduced model. This can be executed, for example, by decomposing the parameter domain Y into K subdomains Y_k and using different affine reduced models for approximating the resulting subsets $\mathcal{M}_k = u(Y_k)$, see, for instance, [6, 15].

1.5. Objective and outline. The standard constructions of reduced models are targeted at making the spaces V_n as efficient as possible for approximating $\mathcal{M}$, that is, making ε_n as small as possible for each given n . For example, for the reduced basis spaces, it is known [2, 12] that a certain greedy selection of snapshots generates spaces V_n such that $\text{dist}(\mathcal{M}, V_n)$ decays at the same rate (polynomial or exponential) as the Kolmogorov n -width

$$(1.25) \quad \delta_n(\mathcal{M}) := \inf\{\text{dist}(\mathcal{M}, E) : \dim(E) = n\}.$$

However these constructions do not ensure the control of μ_n and therefore these reduced spaces may be much less efficient when using the one-space algorithm for the recovery problem.

In view of the above observations, the objective of this paper is to discuss the construction of reduced models (both linear and affine) that are better targeted towards the recovery task. In other words, we want to build the spaces V_n to make the recovery algorithm A_n as efficient as possible, given the measurement space W . Note that a different problem is, given $\mathcal{M}$, to optimize the choice of the measurement functionals ℓ_i picked from some admissible dictionary, which amounts to optimizing the space W , as discussed for example in [4]. Here, we consider our measurement system to be imposed on us, and therefore W to be fixed once and for all.

The rest of our paper is organized as follows. In section 2, we detail the affine map A_n associated with V_n , that can be computed in a similar way as in the linear case. Conversely, we show that any affine recovery map may be interpreted as a one-space algorithm for a certain affine reduced model V_n . For a general set $\mathcal{M}$, the existence and construction of an optimal affine recovery map A_{wca}^* for the worst case error is therefore equivalent to the existence and construction of an optimal reduced space for the recovery problem. We then draw a short comparison with the stochastic setting in which the optimal affine map A_{msa}^* for the mean square error (1.9) is derived explicitly from the second order statistics of u .

In section 3, we compute an approximation of A_{wca}^* by convex optimization, based on a training set of snapshots. Two algorithms are considered: subgradient descent and primal-dual proximal splitting. Our numerical results illustrate the superiority of the latter for this problem. The optimal affine map A_{wca}^* significantly outperforms the one-space algorithm A_{n^*} when standard reduced basis spaces V_n are used and an optimal value n^* is selected using the training set. It also outperforms the affine map A_{msa}^* computed from second order statistics of the training set. All three maps significantly outperform the minimal V -norm recovery given by $A(w) = w = P_W u$.

2. Affine one-space recovery. In this section, we show that any linear recovery algorithm is given by a one-space algorithm and that a similar result holds for any affine algorithm. We then go on to describe the optimal one-space algorithms by exploiting this fact.

2.1. The one-space algorithm. We begin by discussing in more detail the one-space algorithm for a linear space V_n of dimension $n \leq m$. As shown in [3], the map A_n associated with V_n has a simple expression after a proper choice of favorable bases has been made for W and V_n through an SVD applied to the cross-Gramian of an initial pair of orthonormal bases. The resulting *favorable bases* $\{\psi_1, \dots, \psi_m\}$ for W and $\{\varphi_1, \dots, \varphi_n\}$ for V_n satisfy the equations

$$(2.1) \quad \langle \psi_i, \varphi_j \rangle = s_i \delta_{i,j},$$

where

$$(2.2) \quad 1 \geq s_1 \geq s_2 \geq \dots \geq s_n > 0$$

are the singular values of the cross-Gramian. Then, if w is in W , we can write $w = \sum_{j=1}^m w_j \psi_j$ in the favorable basis, and find that

$$(2.3) \quad A_n(w) = \sum_{j=1}^n s_j^{-1} w_j \varphi_j + \sum_{j=n+1}^m w_j \psi_j.$$

Let us observe that the functions ψ_j in the second sum span the space $V_n^\perp \cap W$ while the first sum is the solution of the least-squares problem $\min_{v \in V_n} \|w - P_W v\|$ corrected by the second sum so as to fit the data.

Now consider any linear recovery algorithm $A : W \rightarrow V$. Since we are given the measurement observation w , any algorithm A which is a candidate for optimality must satisfy $P_W(A(w)) = w$ (otherwise the reconstruction error would not be minimized). Thus A should have the form

$$(2.4) \quad A(w) = w + B(w),$$

where $B : W \rightarrow W^\perp$ with $W^\perp$ the orthogonal complement of W in V . Note that in functional analysis the mappings A of the form (2.4) are called liftings.

Therefore, in going further in this paper, we always require that A has the form (2.4) and concentrate on the construction of good linear maps B . Our next observation is that any algorithm A of this form can always be interpreted as a one-space algorithm A_n for a certain space V_n with $n \leq m$.

PROPOSITION 2.1. *Let A be any linear map of the form (2.4). Then, there exists a space V_n of dimension $n \leq m$ such that A coincides with the one-space algorithm (2.3) for V_n .*

Proof. By considering the SVD of the linear transform B , there exists an orthonormal basis $\{\psi_1, \dots, \psi_m\}$ of W and an orthonormal system $\{\omega_1, \dots, \omega_m\}$ in $W^\perp$ such that, with $w = \sum_{j=1}^m w_j \psi_j$,

$$(2.5) \quad Bw = \sum_{j=1}^m \alpha_j w_j \omega_j, \quad w \in W,$$

for some numbers $\alpha_1 \geq \alpha_2 \geq \dots \geq \alpha_m \geq 0$. Defining the functions

$$(2.6) \quad \varphi_j = s_j(\psi_j + \alpha_j \omega_j), \quad s_j = (1 + \alpha_j^2)^{-1/2},$$

and defining V_n as the span of those φ_j for which $\alpha_j \neq 0$, we recover the exact form (2.3) of the one-space algorithm expressed in favorable bases. $\square$

These results can be readily extended to the case where V_n is an affine space given by (1.21) for some given n -dimensional linear space $\tilde{V}_n$ and offset $\bar{u}$. In what follows, we systematically use the notation

$$(2.7) \quad \tilde{u} = u - \bar{u}$$

for the recentered state, and likewise $\tilde{w} = w - \bar{w}$ with $\bar{w} = P_W \bar{u}$ for the recentered observation. The one-space algorithm associated with V_n has the form

$$(2.8) \quad A_n(w) = \bar{u} + \tilde{A}_n(\tilde{w}),$$

where $\tilde{A}_n$ is the one-space linear algorithm associated with $\tilde{V}_n$.

Performances bounds similar to those of the linear case are derived in the same way as in [3]: the reconstruction satisfies

$$(2.9) \quad \|u - A_n(P_W u)\| \leq \mu_n \text{dist}(u, \bar{u} + \tilde{V}_n \oplus (\tilde{V}_n^\perp \cap W)) \leq \mu_n \text{dist}(u, V_n),$$

where

$$(2.10) \quad \mu_n = \mu(\tilde{V}_n, W) = \max_{v \in \tilde{V}_n} \frac{\|v\|}{\|P_W v\|} = s_n^{-1} < \infty.$$

The map A_n is optimal for the cylinders of the form

$$(2.11) \quad \mathcal{K}_n = \{u \in V : \text{dist}(u, V_n) \leq \varepsilon_n\}$$

since it coincides with the Chebyshev ball center of $\mathcal{K}_{n,w} = \mathcal{K}_n \cap V_w$. In particular, one has

$$(2.12) \quad E_{\text{wc}}^*(\mathcal{K}_n) = E_{\text{wc}}(A_n, \mathcal{K}_n) = \mu_n \varepsilon_n.$$

For a solution manifold $\mathcal{M}$ contained in $\mathcal{K}_n$, one has

$$(2.13) \quad E_{\text{wc}}^*(\mathcal{M}) \leq E_{\text{wc}}(A_n, \mathcal{M}) \leq \mu_n \text{dist}(\mathcal{M}, \tilde{V}_n \oplus (\tilde{V}_n^\perp \cap W)) \leq \mu_n \text{dist}(\mathcal{M}, V_n) \leq \mu_n \varepsilon_n,$$

and these inequalities are generally strict.

In view of (2.8) the map A_n is affine. A general affine recovery map takes the form

$$(2.14) \quad A(w) = w + Bw + c,$$

where $B : W \rightarrow W^\perp$ is linear and $c = A(0) \in W^\perp$. The following result is a direct consequence of Proposition 2.1.

COROLLARY 2.2. *Let A be an affine map of the form (2.14). Then, there exists an affine space $V_n = \bar{u} + \tilde{V}_n$ such that A coincides with the one-space algorithm (2.8).*

2.2. The best affine map. In view of this result, the search for an affine reduced model V_n that is best tailored to the recovery problem is equivalent to the search of an optimal affine map. Our next result is that such a map always exist when $\mathcal{M}$ is a bounded set.

THEOREM 2.3. *Let $\mathcal{M}$ be a bounded set. Then there exists a map A_{wca}^* that minimizes $E_{\text{wc}}(A, \mathcal{M})$ among all affine maps A .*

Proof. We consider any affine map A of the form (2.14), so that the error is given by

$$(2.15) \quad E_{\text{wc}}(A, \mathcal{M}) = \sup\{u \in \mathcal{M} : \|P_{W^\perp} u - c - BP_W u\|\} =: F(c, B).$$

We begin by remarking that for each $(c, B) \in W^\perp \times \mathcal{L}(W, W^\perp)$, the map $u \mapsto \|P_{W^\perp} u - c - BP_W u\|$ is uniformly bounded on the bounded set $\mathcal{M}$. Its supremum $F(c, B)$ is thus a finite positive number, which we may write as

$$(2.16) \quad F(c, B) = \sup_{u \in \mathcal{M}} F_u(c, B),$$

where $F_u(c, B) = \|P_{W^\perp} u - c - BP_W u\|$. Each F_u is convex and satisfies the Lipschitz bound

$$(2.17) \quad |F_u(c, B) - F_u(c', B')| \leq \|c - c'\| + M\|B - B'\|_S$$

where $\|B\|_S$

$$(2.18) \quad \|B\|_S = \max\{\|Bv\| : v \in W, \|v\| = 1\}$$

denotes the spectral norm and $M := \sup\{\|P_W u\| : u \in \mathcal{M}\} < \infty$. This implies that the function F is convex and satisfies the same Lipschitz bound.

We note that the linear maps of $\mathcal{L}(W, W^\perp)$ are of rank at most m and therefore, given any orthonormal basis $(e_1, \dots, e_m)$ of W , we can equip $\mathcal{L}(W, W^\perp)$ with the Hilbert–Schmidt norm

$$(2.19) \quad \|B\|_{HS} := \left(\sum_{i=1}^m \|Be_i\|^2 \right)^{1/2},$$

which is equivalent to the spectral norm since

$$(2.20) \quad \|B\|_S \leq \|B\|_{HS} \leq \sqrt{m} \|B\|_S, \quad B \in \mathcal{L}(W, W^\perp).$$

In particular F is continuous with respect to the Hilbertian norm

$$(2.21) \quad \|(c, B)\|_H := \left(\|c\|^2 + \sum_{i=1}^m \|Be_i\|^2 \right)^{1/2}.$$

The function F may not be infinite at infinity: this happens if there exists a nontrivial pair (c, B) such that

$$c + BP_W u = 0, \quad u \in \mathcal{M}.$$

In order to fix this problem, we define the subspace

$$(2.22) \quad S_0 := \left\{ (c, B) \in W^\perp \times \mathcal{L}(W, W^\perp) : c + BP_W u = 0, u \in \mathcal{M} \right\}.$$

and we denote by S_1 its orthogonal complement in $W^\perp \times \mathcal{L}(W, W^\perp)$ for the inner product associated with the above Hilbertian norm $\|\cdot\|_H$. The function F is constant in the direction of S_0 and therefore we are left to prove the existence of the minimum of F on S_1 . For any $(c, B) \in S_1$, there exists $u \in \mathcal{M}$ such that $c + BP_W u \neq 0$. This implies that

$$(2.23) \quad \lim_{|t| \rightarrow +\infty} \|P_{W^\perp} u - tc - tBP_W u\| = +\infty$$

and, therefore, that $\lim_{|t| \rightarrow +\infty} F_u(t(c, B)) = +\infty$. This shows that F is infinite at infinity when restricted to S_1 . Any convex and continuous function in a Hilbert space is weakly lower semicontinuous, and admits a minimum when it is infinite at infinity. We thus conclude the existence of a minimizer (c^*, B^*) of F and therefore

$$(2.24) \quad A_{\text{wca}}^*(w) = w + c^* + B^* w$$

is an optimal affine recovery map. $\square$

2.3. The best affine map in the stochastic setting. In the stochastic setting, assuming that u has finite second order moments, the optimal map that minimizes the mean square error (1.9) is given by the conditional expectation

$$(2.25) \quad A_{\text{ms}}^*(w) = \mathbb{E}(u \mid P_W u = w),$$

that is, the expectation of the posterior distribution p_w of u conditioned to the observation of w . Various sampling strategies have been developed in order to approximate the posterior and its expectation; see [13] for a survey. These approaches come at a significant computational cost since they require a specific sampling for each instance w of observed data. In the parametric PDE setting, each sample requires one solve of the forward problem.

On the other hand, it is well known that an optimal affine map A_{msa}^* for the mean square error can be explicitly derived from the first and second order statistics of u . We briefly recall this derivation by using an arbitrary orthonormal basis $(e_1, \dots, e_m)$ of W that we complement into an orthonormal basis $(e_j)_{j \geq 1}$ of V . We write

$$(2.26) \quad u = \sum_{j \geq 1} w_j e_j \quad \text{and} \quad \bar{u} = \mathbb{E}(u) = \sum_{j \geq 1} \bar{w}_j e_j, \quad \bar{w}_j := \mathbb{E}(w_j),$$

as well as

$$(2.27) \quad \tilde{u} = u - \bar{u} = \sum_{j \geq 1} \tilde{w}_j e_j, \quad \tilde{w}_j := w_j - \bar{w}_j.$$

An affine recovery map of the form (2.14) leaves the coordinates $w_1, \dots, w_m$ unchanged and recovers for each $i \geq 1$

$$(2.28) \quad w_{m+i}^* = c_i + \sum_{j=1}^m b_{i,j} w_j,$$

which can be rewritten as

$$(2.29) \quad w_{m+i}^* = \bar{w}_{m+i} + d_i + \sum_{j=1}^m b_{i,j} \tilde{w}_j.$$

Since $E_{\text{ms}}(A) = \sum_{i \geq 1} \mathbb{E}(|w_{m+i}^* - w_{m+i}|^2)$, the numbers d_i and $b_{i,j}$ are found by separately minimizing each term. By the Pythagorean theorem one has

$$(2.30) \quad \mathbb{E}(|w_{m+i}^* - w_{m+i}|^2) = |d_i|^2 + \mathbb{E}\left(\left|\sum_{j=1}^m b_{i,j} \tilde{w}_j - \tilde{w}_{m+i}\right|^2\right),$$

which shows that we should take $d_i = 0$. Minimizing the second term leads to the orthogonal projection equations

$$(2.31) \quad \sum_{j=1}^m b_{i,j} t_{j,l} = t_{m+j,l}, \quad l = 1, \dots, m,$$

which involve the entries of the covariance matrix

$$(2.32) \quad \mathbf{S} := (t_{i,j}), \quad t_{i,j} := \mathbb{E}(\tilde{w}_i \tilde{w}_j).$$

Therefore, with the block decomposition

$$(2.33) \quad \mathbf{S} = \begin{pmatrix} \mathbf{S}_{1,1} & \mathbf{S}_{1,2} \\ \mathbf{S}_{2,1} & \mathbf{S}_{2,2} \end{pmatrix},$$

corresponding to the splitting of rows and columns from $\{1, \dots, m\}$ and $\{m+1, m+2, \dots\}$, one obtains that the matrix $\mathbf{B} = (b_{i,j})$ that defines the optimal affine map satisfies $\mathbf{S}_{1,1} \mathbf{B}^T = \mathbf{S}_{1,2}$ and, therefore,

$$(2.34) \quad \mathbf{B} = \mathbf{S}_{2,1} \mathbf{S}_{1,1}^{-1},$$

where we have used the symmetry of $\mathbf{S}$. In other words,

$$(2.35) \quad A_{\text{msa}}^*(w) = w + P_{W^\perp} \bar{u} + B \tilde{w},$$

where the linear transform $B \in \mathcal{L}(W, W^\perp)$ is represented by the matrix $\mathbf{B}$ in the basis $(e_j)_{j \geq 1}$.

The optimal affine recovery map A_{msa}^* agrees with the optimal map A_{ms}^* in the particular case where u has Gaussian distribution, therefore entirely characterized by its average $\bar{u}$ and covariance matrix $\mathbf{S}$. To see this, assume for simplicity that V is finite dimensional. The distribution of $\mathbf{u} = (w_j)_{j \geq 1}$ has density proportional to $\exp(-\frac{1}{2}\langle \mathbf{T}\tilde{\mathbf{u}}, \tilde{\mathbf{u}} \rangle)$, where $\mathbf{T} = \mathbf{S}^{-1}$. We expand the quadratic form into

$$(2.36) \quad \frac{1}{2}\langle \mathbf{T}\tilde{\mathbf{u}}, \tilde{\mathbf{u}} \rangle = \frac{1}{2}\langle \mathbf{T}_{1,1}\tilde{\mathbf{w}}, \tilde{\mathbf{w}} \rangle + \langle \mathbf{T}_{2,1}\tilde{\mathbf{w}}, \tilde{\mathbf{w}}_\perp \rangle + \frac{1}{2}\langle \mathbf{T}_{2,2}\tilde{\mathbf{w}}_\perp, \tilde{\mathbf{w}}_\perp \rangle,$$

where $\tilde{\mathbf{w}}_\perp = (\tilde{w}_{m+j})_{j \geq 1}$ and $\tilde{\mathbf{w}} = (\tilde{w}_j)_{j=1,\dots,m}$, and where

$$(2.37) \quad \mathbf{T} = \begin{pmatrix} \mathbf{T}_{1,1} & \mathbf{T}_{1,2} \\ \mathbf{T}_{2,1} & \mathbf{T}_{2,2} \end{pmatrix}$$

is a block decomposition similar to that of $\mathbf{S}$. The distribution of the vector $\tilde{\mathbf{w}}_\perp$ conditional to the observation of $\tilde{\mathbf{w}}$ is also Gaussian and its expectation coincides with the minimum of the quadratic form

$$(2.38) \quad Q_{\mathbf{w}}(\tilde{\mathbf{w}}_\perp) = \frac{1}{2}\langle \mathbf{T}_{2,2}\tilde{\mathbf{w}}_\perp, \tilde{\mathbf{w}}_\perp \rangle + \langle \mathbf{T}_{2,1}\tilde{\mathbf{w}}, \tilde{\mathbf{w}}_\perp \rangle.$$

Therefore

$$(2.39) \quad \mathbb{E}(\tilde{\mathbf{w}}_\perp | \tilde{\mathbf{w}}) = -\mathbf{T}_{2,2}^{-1}\mathbf{T}_{2,1}\tilde{\mathbf{w}} = \mathbf{S}_{2,1}\mathbf{S}_{1,1}^{-1}\tilde{\mathbf{w}} = \mathbf{B}\tilde{\mathbf{w}},$$

which shows that

$$(2.40) \quad A_{\text{ms}}^*(w) = \mathbb{E}(u | P_W u = w) = A_{\text{msa}}^*(w).$$

One main interest of the above discussed stochastic setting is that the best affine map is now explicitly given by the second order statistics, in view of (2.34). This contrasts with the deterministic setting in which the optimal affine map is obtained by minimization of the convex functional F from (2.16) and does not generally have a simple explicit expression. Algorithms for solving this minimization problem are the object of the next section.

Only for particular cases where $\mathcal{M}$ has a simple geometry does the best affine map A_{wca}^* in the deterministic setting have a simple expression. One typical example is when $\mathcal{M}$ is an ellipsoid described by an equation of the form

$$(2.41) \quad \langle \mathbf{T}\tilde{\mathbf{u}}, \tilde{\mathbf{u}} \rangle \leq 1$$

for a symmetric positive matrix $\mathbf{T}$. Then, the set $\mathcal{M}_w = \mathcal{M} \cap V_w$ is also an ellipsoid associated with the above quadratic form $Q_{\mathbf{w}}$. The coordinates of its center are therefore given by the equation $\tilde{\mathbf{w}}_\perp = -\mathbf{T}_{2,2}^{-1}\mathbf{T}_{2,1}\tilde{\mathbf{w}}$, which is the same as that defining the conditional expectation in (2.39). This shows that, in the particular case of an ellipsoid, (i) the optimal map A_{wc}^* agrees with the optimal affine recovery map A_{wca}^* for the worst case error, and (ii) it has an explicit expression which agrees with the optimal map A_{msa}^* for the mean square error when the prior is a Gaussian with $\mathbf{T}$ as inverse covariance matrix.

3. Algorithms for optimal affine recovery.

3.1. Discretization and truncation. We have seen that the optimal affine recovery map is obtained by minimizing the convex function

$$(3.1) \quad F(c, B) = \sup_{u \in \mathcal{M}} \|P_{W^\perp} u - c - BP_W u\|,$$

over $W^\perp \times \mathcal{L}(W, W^\perp)$. This optimization problem cannot be solved exactly for two reasons:

- (i) The sets $W^\perp$ as well as $\mathcal{L}(W, W^\perp)$ are infinite dimensional when V is infinite dimensional.
- (ii) One single evaluation of $F(c, B)$ requires, in principle, to explore the entire manifold $\mathcal{M}$.

The first difficulty is solved by replacing V by a subspace Z_N of finite dimension $\dim(Z_N) = N$ that approximates the solution manifold $\mathcal{M}$ with an accuracy of smaller order than that expected for the recovery error. One possibility is to use a finite element space $Z_N = V_h$ of sufficiently small mesh size h . However its resulting dimension $N = N(h)$ needed to reach the accuracy could still be quite large. An alternative is to use reduced model spaces Z_N which are more efficient for the approximation of $\mathcal{M}$, as we discuss further.

We therefore minimize $F(c, B)$ over $\widetilde{W}^\perp \times \mathcal{L}(W, \widetilde{W}^\perp)$, where $\widetilde{W}^\perp$ is the orthogonal complement of W in the space $W + Z_N$, and obtain an affine map $\tilde{A}_{\text{wca}}$ defined by

$$(3.2) \quad \tilde{A}_{\text{wca}}(w) = w + \tilde{c} + \tilde{B}w, \quad (\tilde{c}, \tilde{B}) := \operatorname{argmin}\{F(c, B) : c \in \widetilde{W}^\perp, B \in \mathcal{L}(W, \widetilde{W}^\perp)\}.$$

In order to compare its performance with that of A_{wca}^* , we first observe that

$$(3.3) \quad \|P_{W^\perp} u - P_{\widetilde{W}^\perp} u\| \leq \varepsilon_N := \sup_{u \in \mathcal{M}} \operatorname{dist}(u, Z_N).$$

For any $(c, B) \in W^\perp \times \mathcal{L}(W, W^\perp)$, we define $(\tilde{c}, \tilde{B}) \in \widetilde{W}^\perp \times \mathcal{L}(W, \widetilde{W}^\perp)$ by $\tilde{c} = P_{\widetilde{W}^\perp} c$ and $\tilde{B} = P_{\widetilde{W}^\perp} \circ B$. Then, for any $u \in \mathcal{M}$,

$$\begin{aligned} \|P_{W^\perp} u - \tilde{c} - \tilde{B}u\| &\leq \|P_{\widetilde{W}^\perp}(P_{W^\perp} u - c - BP_W u)\| + \|P_{W^\perp} u - P_{\widetilde{W}^\perp} u\| \\ &\leq \|P_{W^\perp} u - c - BP_W u\| + \varepsilon_N. \end{aligned}$$

It follows that we have the framing

$$(3.4) \quad E(A_{\text{wca}}^*, \mathcal{M}) \leq E(\tilde{A}_{\text{wca}}, \mathcal{M}) \leq E(A_{\text{wca}}^*, \mathcal{M}) + \varepsilon_N,$$

which shows that the loss in the recovery error is at most of the order ε_N .

To understand how large N should be, let us observe that a recovery map A of the form (2.14) takes its value in the linear space

$$(3.5) \quad F_{m+1} = \mathbb{R}c + \operatorname{ran}(B),$$

which has dimension $m+1$. It follows that the recovery error is always larger than the approximation error by such a space. Therefore

$$(3.6) \quad E_{\text{wc}}(A_{\text{wca}}^*, \mathcal{M}) \geq \delta_{m+1}(\mathcal{M}),$$

where $\delta_{m+1}(\mathcal{M})$ is the Kolmogorov n -width defined by (1.25) for $n = m+1$. Therefore, if we could use the space $Z_n := E_n$ that exactly achieve the infimum in (1.25), we

would be ensured that, with $N = m + 1$, the additional error $\varepsilon_N = \delta_{m+1}(\mathcal{M})$ in (3.4) is of smaller order than $E_{wc}(A_{wca}^*, \mathcal{M})$. As a result we would obtain the framing

$$(3.7) \quad E(A_{wca}^*, \mathcal{M}) \leq E(\tilde{A}_{wca}, \mathcal{M}) \leq 2E(A_{wca}^*, \mathcal{M}).$$

In practice, since we do not have access to the n -width spaces, we use instead the reduced basis spaces $Z_n := V_n$ which are expected to have comparable approximation performances in view of the results from [2, 12]. We take N larger than m but of comparable order.

The second difficulty is solved by replacing the set $\mathcal{M}$ in the supremum that defines $F(c, B)$ by a discrete training set $\tilde{\mathcal{M}}$, which corresponds to a discretization $\tilde{Y}$ of the parameter domain Y , that is,

$$(3.8) \quad \tilde{\mathcal{M}} := \{u(y) : y \in \tilde{Y}\}$$

with finite cardinality.

We therefore minimize over $\tilde{W}^\perp \times \mathcal{L}(W, \tilde{W}^\perp)$ the function

$$(3.9) \quad \tilde{F}(c, B) = \sup_{u \in \tilde{\mathcal{M}}} \|P_{W^\perp} u - c - BP_W u\|,$$

which is computable. The additional error resulting from this discretization can be controlled from the resolution of the discretization. Namely, let $\varepsilon > 0$ be the smallest value such that $\tilde{\mathcal{M}}$ is an ε -approximation net of $\mathcal{M}$, that is, $\mathcal{M}$ is covered by the V -balls $B(u, \varepsilon)$ for $u \in \tilde{\mathcal{M}}$. Then, we find that

$$(3.10) \quad \tilde{F}(c, B) \leq F(c, B) \leq \tilde{F}(c, B) + \varepsilon \|B\|_{\mathcal{L}(W, \tilde{W}^\perp)},$$

which shows that the additional recovery error will be of the order of ε amplified by the norm of the linear part of the optimal recovery map.

One difficulty is that the cardinality of ε -approximation nets become potentially intractable for small ε as the parameter dimension becomes large, due to the curse of dimensionality. This difficulty also occurs when constructing a reduced basis by a greedy selection process which also needs to be performed in sufficiently dense discretized sets. Recent results obtained in [9] show that in certain relevant instances ε approximation nets can be replaced by random training sets of smaller cardinality. One interesting direction for further research is to apply similar ideas in the context of the present paper.

3.2. Optimization algorithms. As already brought up in the previous section, the practical computation of $\tilde{A}_{wc}$ consists in solving

$$(3.11) \quad \min_{(c, B) \in \tilde{W}^\perp \times \mathcal{L}(W, \tilde{W}^\perp)} \sup_{u \in \tilde{\mathcal{M}}} \|P_{W^\perp} u - c - BP_W u\|^2.$$

The numerical solution of this problem is challenging due to its lack of smoothness (the objective function is convex but nondifferentiable) and its high dimensionality (for a given target accuracy ε_N , the cardinality of $\tilde{\mathcal{M}}$ might be large). One could use classical subgradient methods, which are simple to implement. However these schemes only guarantee a very slow $O(k^{-1/2})$ convergence rate of the objective function, where k is the number of iterations. This approach did not give satisfactory results in our case; due to the slow convergence, the solution update of one iteration falls below machine precision before approaching the minimum close enough; see Figure 3.1. This

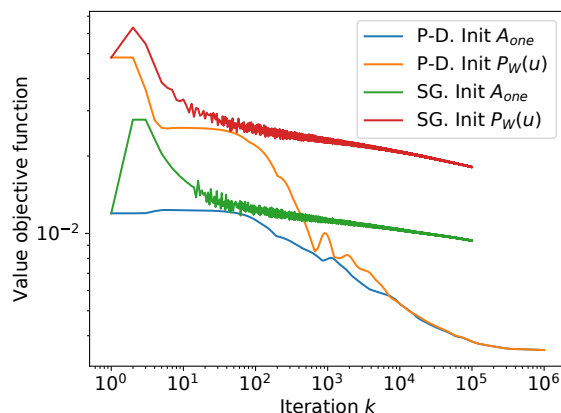


FIG. 3.1. Convergence of the objective function for two different optimization algorithms and starting guesses. P.D. = Primal-Dual splitting. S.G. = Subgradient. Here, $m = 40$.

has motivated the use of a primal-dual splitting method which is known to ensure a $O(1/k)$ convergence rate on the partial duality gap. We next describe this method, but only briefly, as a detailed analysis would make us deviate too far from the main topic of this paper. A complete analysis with further examples of application will be presented in a forthcoming work [14].

We assume without loss of generality that $\dim(W + V_N) = m + N$ and that $\dim \widetilde{W}^\perp = N$. Let $\{\psi_i\}_{i=1}^{m+N}$ be an orthonormal basis of $W + V_N$ such that $W = \text{span}\{\psi_1, \dots, \psi_m\}$. Since for any $u \in V$,

$$P_{W+V_N}u = \sum_{i=1}^{m+N} u_i \psi_i,$$

the components of u in W can be given in terms of the vector $\mathbf{w} = (u_i)_{i=1}^m$ and the ones in $\widetilde{W}^\perp$ with $\mathbf{u} = (u_{i+m})_{i=1}^N$.

We now consider the finite training set

$$(3.12) \quad \widetilde{\mathcal{M}} := \{u^1, \dots, u^J\}, \quad J := \#(\widetilde{\mathcal{M}}) < \infty,$$

and denote by $\mathbf{w}^j$ and $\mathbf{u}^j$ the vectors associated with the snapshot functions u^j for $j = 1, \dots, J$. One may express the problem (3.11) as the search for

$$(3.13) \quad \min_{(\mathbf{R}, \mathbf{b}) \in \mathbb{R}^{N \times m} \times \mathbb{R}^N} \max_{j=1, \dots, J} \|\mathbf{u}^j - \mathbf{R}\mathbf{w}^j - \mathbf{b}\|_2^2.$$

Concatenating the matrix and vector variables $(\mathbf{R}, \mathbf{b})$ into a single $\mathbf{x} \in \mathbb{R}^{m(N+1)}$, we rewrite the above problem as

$$(3.14) \quad \min_{\mathbf{x} \in \mathbb{R}^{m(N+1)}} \max_{j=1, \dots, J} f_j(\mathbf{Q}_j \mathbf{x}),$$

where $\mathbf{Q}_j \in \mathbb{R}^{N \times m(N+1)}$ is a sparse matrix built using the coefficients of $\mathbf{w}^j$ and $f_j(\mathbf{y}) := \|\mathbf{u}^j - \mathbf{y}\|_2^2$.

The key observation to build our algorithm is that problem (3.14) can be equivalently written as a minimization problem on the epigraphs, i.e.,

$$(3.15) \quad \begin{aligned} & \min_{(\mathbf{x}, t) \in \mathbb{R}^{m(N+1)} \times \mathbb{R}^+} t \quad \text{subject to} \quad f_j(\mathbf{Q}_j \mathbf{x}) \leq t, \quad j = 1, \dots, J \\ \iff & \min_{(\mathbf{x}, t) \in \mathbb{R}^{m(N+1)} \times \mathbb{R}^+} t \quad \text{subject to} \quad (\mathbf{Q}_j \mathbf{x}, t) \in \text{epi}_{f_j}, \quad j = 1, \dots, J, \end{aligned}$$

or, in a more compact (and implicit) form,

$$(P_{\text{epi}}) \quad \min_{(\mathbf{x}, t) \in \mathbb{R}^{m(N+1)} \times \mathbb{R}^+} t + \sum_{j=1}^J \iota_{\text{epi}_{f_j}}(\mathbf{Q}_j \mathbf{x}, t),$$

where, for any nonempty set S the indicator function ι_S has value 0 on S and $+\infty$ on S^c .

This problem takes the following canonical expression, which is amenable to a primal-dual proximal splitting algorithm:

$$(3.16) \quad \min_{(\mathbf{x}, t) \in \mathbb{R}^{m(N+1)} \times \mathbb{R}} G(\mathbf{x}, t) + F \circ L(\mathbf{x}, t).$$

Here, G is the projection map for the second variable

$$(3.17) \quad G(\mathbf{x}, t) = t,$$

the linear operator L is defined by

$$(3.18) \quad L(\mathbf{x}, t) := ((\mathbf{Q}_1 \mathbf{x}, t), (\mathbf{Q}_2 \mathbf{x}, t), \dots, (\mathbf{Q}_J \mathbf{x}, t))$$

and acts from $\mathbb{R}^{m(N+1)} \times \mathbb{R}$ to $\times_{j=1}^J (\mathbb{R}^N \times \mathbb{R})$ and the function F acting from $\times_{j=1}^J (\mathbb{R}^N \times \mathbb{R})$ to $\mathbb{R}$ is defined by

$$(3.19) \quad F\left((\mathbf{v}_1, t_1), \dots, (\mathbf{v}_J, t_J)\right) := \sum_{j=1}^J \iota_{\text{epi}_{f_j}}(\mathbf{v}_j, t_j).$$

Note that F is the indicator function of the Cartesian product of epigraphs.

Before introducing the primal-dual algorithm, some remarks are in order:

- (i) We recall that if ϕ is a proper closed convex function on $\mathbb{R}^d$, its proximal mapping prox_ϕ is defined by

$$(3.20) \quad \text{prox}_\phi(y) = \underset{x \in \mathbb{R}^d}{\text{argmin}} \left(\phi(x) + \frac{1}{2} \|x - y\|_2^2 \right).$$

- (ii) The adjoint operator L^* is given by

$$(3.21) \quad L^*\left((\mathbf{v}_1, t_1), \dots, (\mathbf{v}_J, t_J)\right) := \left(\sum_{j=1}^J \mathbf{Q}_j^T \mathbf{v}_j, \sum_{j=1}^J t_j \right).$$

It can be easily shown that the operator norm of L satisfies $\|L\|^2 \leq J + \sum_{j=1}^J \|\mathbf{Q}_j\|^2$.

- (iii) Both G and F are simple functions in the sense that their proximal mappings, prox_G and prox_F , can be computed in closed form. See [14] for details.

The iterations of our primal-dual splitting method read for $k \geq 0$,
(3.22)

$$\begin{aligned}(\mathbf{x}, t)^{k+1} &= \text{prox}_{\gamma_G G} \left((\mathbf{x}, t)^k - \gamma_G L^* \left(\left((\mathbf{v}_1, \xi_1), \dots, (\mathbf{v}_J, \xi_J) \right)^k \right) \right), \\ (\bar{\mathbf{x}}, \bar{t})^{k+1} &= (\mathbf{x}, t)^{k+1} + \theta \left((\mathbf{x}, t)^{k+1} - (\mathbf{x}, t)^k \right), \\ \left((\mathbf{v}_1, \xi_1), \dots, (\mathbf{v}_J, \xi_J) \right)^{k+1} &= \text{prox}_{\gamma_F \hat{F}} \left(\left((\mathbf{v}_1, \xi_1), \dots, (\mathbf{v}_J, \xi_J) \right)^k + \gamma_F L(\bar{\mathbf{x}}, \bar{t})^{k+1} \right),\end{aligned}$$

where $\hat{F}$ is the Fenchel–Legendre transform of F , $\gamma_G > 0$ and $\gamma_F > 0$ are such that $\gamma_G \gamma_F < 1/\|L\|^2$, and $\theta \in [-1, +\infty[$ (it is generally set to $\theta = 1$ as in [8]).

Algorithm 1 gives some guidelines and summarizes in an informal pseudocode style the main steps of the primal-dual approach (the implementation of the routine “BuildQ” is left to the reader).

Algorithm 1 Primal-dual algorithm: $\mathbf{R}, \mathbf{b} = \text{PD}(\widetilde{\mathcal{M}}, \mathcal{M}_{\text{greedy}}, W, K_{\max})$.

```

1: Input:
    • training manifold  $\widetilde{\mathcal{M}}$  for primal-dual iterations
    • training manifold  $\mathcal{M}_{\text{greedy}}$  for greedy algorithm
    • basis  $\{\omega_i\}_{i=1}^m$  of measurement space  $W$ 
    • maximum number of iteration  $K_{\max}$ 
2: Generate basis  $\{v_i\}_{i=1}^N$  of  $V_N$  // e.g., with a greedy algorithm over  $\mathcal{M}_{\text{greedy}}$ ;
   see (3.27)
3: Build orthonormal basis  $\{\psi_i\}_{i=1}^{m+N}$  of  $W + V_N$  with a Gram–Schmidt procedure
   over  $\{w_1, \dots, w_m, v_1, \dots, v_N\}$ . In this way,  $\widetilde{W}^\perp = \text{span}\{\psi_{m+1}, \dots, \psi_{m+N}\}$ 
4: Qlist, wlist, ulist = []
5: for all  $u \in \widetilde{\mathcal{M}}$  do // Build matrices  $\mathbf{Q}^j$  of (3.14)
6:    $\mathbf{w} = \{\langle u, \psi_i \rangle\}_{i=1}^m$ ,  $\mathbf{u} = \{\langle u, \psi_i \rangle\}_{i=m+1}^{m+N}$ ,  $\mathbf{Q} = \text{BuildQ}(\mathbf{w})$ 
7:   Qlist.append( $\mathbf{Q}$ ), wlist.append( $\mathbf{w}$ ), ulist.append( $\mathbf{u}$ )
8: end for
9: Estimate  $\|L\|$  // e.g., with power method
10: Set  $\gamma_G$  and  $\gamma_F$  such that  $\gamma_G \gamma_F < 1/\|L\|^2$ 
11:  $\bar{\mathbf{x}} = \mathbf{x} = \text{zeros}(m(N+1))$  // starting guess for  $A : W \rightarrow V$  set to  $A(w) = w$ 
12:  $t = 1$  // starting guess for  $t > 0$ 
13:  $\left( (\mathbf{v}_1, \xi_1), \dots, (\mathbf{v}_J, \xi_J) \right) = L(\mathbf{x}, t)$  // starting guess dual variables
14: for  $k$  in  $[0, K_{\max}]$  do // primal-dual iterations
15:    $(\mathbf{x}_{\text{old}}, t_{\text{old}}) = (\mathbf{x}, t)$ 
16:    $\left( (\mathbf{v}_1, \xi_1), \dots, (\mathbf{v}_J, \xi_J) \right) = \text{prox}_{\gamma_F \hat{F}} \left( \left( (\mathbf{v}_1, \xi_1), \dots, (\mathbf{v}_J, \xi_J) \right) + \gamma_F L(\bar{\mathbf{x}}, \bar{t}) \right)$ 
17:    $(\mathbf{x}, t) = \text{prox}_{\gamma_G G} \left( (\mathbf{x}, t) - \gamma_G L^* \left( \left( (\mathbf{v}_1, \xi_1), \dots, (\mathbf{v}_J, \xi_J) \right) \right) \right)$ 
18:    $(\bar{\mathbf{x}}, \bar{t}) = (\mathbf{x}, t) + \theta \left( (\mathbf{x}, t) - (\mathbf{x}_{\text{old}}, t_{\text{old}}) \right)$ 
19: end for
20: Retrieve  $\mathbf{R}, \mathbf{c}$  by appropriately reshaping  $\mathbf{x}$ 
21: Output:  $\mathbf{R}, \mathbf{c}$ 

```

To illustrate the relevance of this algorithm for our purposes, we compare its performance with a standard subgradient method. Figure 3.1 plots the convergence history of the objective function across the iterations of both optimization methods

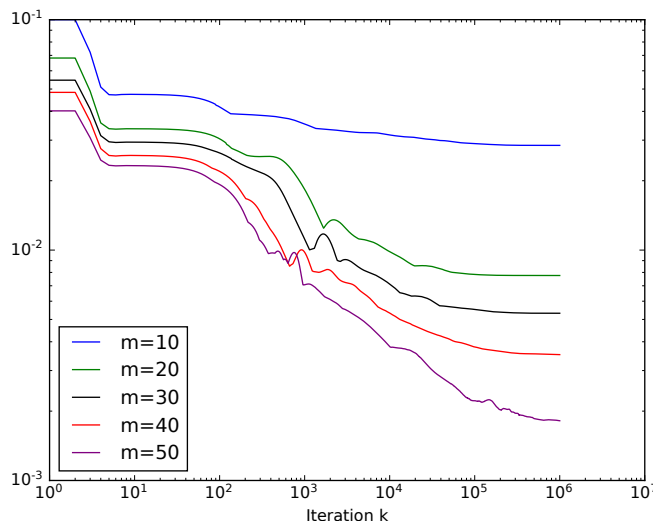


FIG. 3.2. Convergence of the objective function in the primal-dual iterations for $m = 10, 20, 30, 40, 50$.

for the example described in the next section ($m = 40$, $N = 110$, and $J = 10^3$). Two different reconstruction maps have been considered as starting guesses: the minimal V -norm recovery map given by $A(w) = w = P_W u$, and the one-space algorithm A_n^* based on reduced basis spaces V_n with an optimal choice n^* for n . The convergence plot shows the superiority of the primal-dual method which converges to the same minimal value of the objective function after 10^5 iterations regardless of the initialization, while the subgradient method fails to reach it since its increments fall below machine precision.

For the same numerical example described next, we vary m and consider $m = 10, 20, 30, 40, 50$. Figure 3.2 gives the convergence of the reconstruction error over the training set $\mathcal{M}$ across the primal-dual iterations (for simplicity, we took P_{W_m} as the starting guess for $A_{wca}^{(m)}$). To make sure that we reach convergence, we perform 10^6 iterations for each case. As expected, we observe in this figure that the final value of the objective function decreases as we increase the value of m (the reconstruction error decreases as we increase the number of measurements).

3.3. Numerical tests. We present some numerical experiments, aiming primarily at comparing the three above discussed recovery maps in terms of the maximum reconstruction error: the one-space affine map A_n , the best affine map A_{msa}^* for the mean square error, and the best affine map A_{wca}^* for the worst case error. In addition, we also consider the minimum V norm reconstruction map $A(w) = w = P_W u$. The results highlight the superiority of the best affine algorithm with respect to the reconstruction error. This comes however at the cost of a computationally intensive training phase as previously described.

We consider the elliptic problem

$$(3.23) \quad \begin{cases} -\operatorname{div}(a(y)\nabla u) = f, & x \in D, \\ u(x) = 0, & x \in \partial D, \end{cases}$$

on the unit square $D =]0, 1[^2$ with a certain parameter dependence in the field a .

More precisely, for a given $p \geq 1$, we consider “checkerboard” random fields where $a(y)$ is piecewise constant on a $p \times p$ subdivision of the unit square:

$$D = \bigcup_{i,j=0}^{p-1} S_{i,j}$$

with

$$S_{i,j} := \left[\frac{i}{p}, \frac{i+1}{p} \right] \times \left[\frac{j}{p}, \frac{j+1}{p} \right], \quad i, j \in 0, \dots, p-1.$$

The random field is defined as

$$(3.24) \quad a(y) = 1 + \frac{1}{2} \sum_{i,j=0}^{p-1} \chi_{S_{i,j}} y_{i,j},$$

where χ_S denotes the characteristic function of a set S , and the $y_{i,j}$ are random coefficients that are independent each with identical uniform distribution on $[-1, 1]$. Thus, our vector of parameters is

$$\mathbf{y} = (y_{i,j})_{i,j=0}^{p-1} \in \mathbb{R}^{p \times p}.$$

In our numerical tests, we take $p = 4$, that is, 16 parameters, and work in the ambient space $V = H_0^1(D)$. All the sets of snapshots used for training and validating the reconstruction algorithms have been computed by first generating a certain number J of random parameters $\mathbf{y}^1, \dots, \mathbf{y}^J$ with each $\mathbf{y}^i \in [-1, 1]^{p \times p}$, and then solving the variational form of (3.23) in $V = H_0^1(D)$ using $\mathbb{P}_1$ finite elements on a regular grid of mesh size $h = 2^{-7}$. This gives the corresponding solutions $u_h^i = u_h(\mathbf{y}^i)$ that are used in the computations. To ease the reading, in the following we drop the dependence on h in the notation.

The sensor measurements are modeled with linear functionals that are local averages of the form

$$(3.25) \quad \ell_{\mathbf{x},\tau}(u) = \int_D u(\mathbf{r}) \varphi_\tau(\mathbf{r} - \mathbf{x}) \, d\mathbf{r},$$

where

$$(3.26) \quad \varphi_\tau(\mathbf{r}) \propto \exp(-|\mathbf{r}|/2\tau^2)$$

is a radial function such that $\int \varphi_\tau = 1$. The parameter $\tau > 0$ represents the spread around the center $\mathbf{x}$. For the observation space W of our example, we randomly select $m = 50$ centers $\mathbf{x}_i \in [0.1, 0.9]^2$ and spreads $\tau_i \in [0.05, 0.1]$, and compute the Riesz representers $\omega_{\mathbf{x}_i,\tau}$ of $\ell_{\mathbf{x}_i,\tau}$ in $H_0^1(D)$. We then set

$$W := \text{span}\{\omega_{\mathbf{x}_i,\tau}\}_{i=1}^M$$

which is a space of dimension $m = 50$. Figure 3.3 shows the m centers $\mathbf{x}_i$. As an example, the figure also plots the function $\omega_{\mathbf{x}_i,\tau}$ for $i = 10$, which has center $\mathbf{x}_i = (0.23, 0.75)$ and spread $\tau_i = 0.06$.

As explained in section 3.1, the first step to compute the best algorithm in practice consists in replacing $V = H_0^1(D)$ by a finite dimensional space that approximates the solution manifold $\mathcal{M}$ at an accuracy smaller than the one expected for the recovery

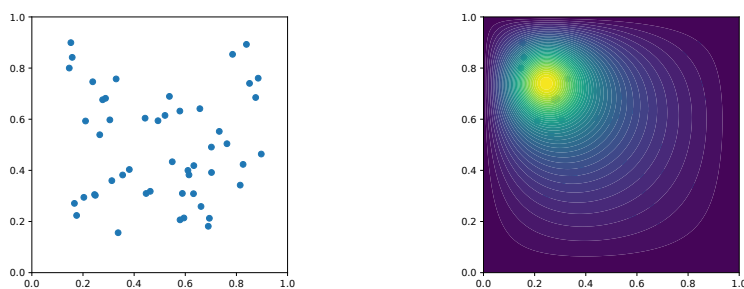


FIG. 3.3. Sensor locations and the function $\omega_{\mathbf{x}_i, \tau_i}$ for $i = 10$ ($\mathbf{x}_i = (0.23, 0.75)$ and $\tau_i = 0.06$).

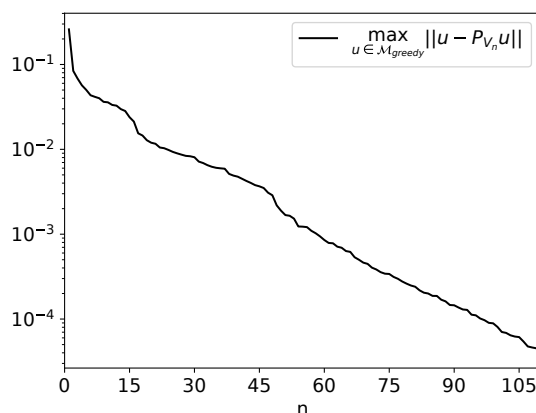


FIG. 3.4. Greedy algorithm: decay of the error $e_n^{(\text{greedy})} = \max_{u \in \mathcal{M}_{\text{greedy}}} \|u - P_{V_n} u\|$.

error. Here, we replace V by $W + V_N$, where V_N is a reduced basis of dimension $N = 110$ that has been generated by running the classical greedy algorithm from [7] over a training set $\mathcal{M}_{\text{greedy}}$ of 10^3 snapshots. We recall that an idealized version is defined for $n \geq 1$ as

$$(3.27) \quad u_n \in \operatorname{argmax}_{u \in \mathcal{M}_{\text{greedy}}} \|u - P_{V_{n-1}} u\|, \quad V_n := V_{n-1} \oplus \mathbb{R}u_n = \operatorname{span}\{u_1, \dots, u_n\},$$

with the convention $V_0 := \{0\}$. Figure 3.4 gives the decay of the error

$$e_n^{(\text{greedy})} = \max_{u \in \mathcal{M}_{\text{greedy}}} \|u - P_{V_n} u\|$$

across the greedy iterations.

We next estimate the truncation accuracy ε_N defined in (3.3). This has been done by computing the maximum of the error $\|u - P_{V_N} u\|$ over the training set $\mathcal{M}_{\text{greedy}}$ supplemented by a test set $\mathcal{M}_{\text{test}}$, also of 10^3 snapshots. We obtain the estimate

$$\max_{u \in \mathcal{M}_{\text{greedy}} \cup \mathcal{M}_{\text{test}}} \|u - P_{V_N} u\| \leq \varepsilon_N = 5.10^{-5}.$$

In the comparison of the three different reconstruction algorithms, we want to illustrate the impact of the number of measurements that are used. To do this, we consider the nested subspaces

$$W_m = \operatorname{span}\{\omega_{\mathbf{x}_i, \tau_i}\}_{i=1}^m \subset W$$

for $m = 10, 20, 30, 40, 50$ so that $W_{50} = W$.

For the computation of the best affine algorithm, we generate a new training set $\widetilde{\mathcal{M}}$ of 10^3 snapshots which we project into $W + V_N$. This projected set, which we denote by $P_{W+V_N}\widetilde{\mathcal{M}}$ with a slight abuse of notation, is used to compute

$$\widetilde{A}_{\text{wca}}^{(m)}(u) = \widetilde{c}^{(m)} + \widetilde{B}^{(m)}P_{W_m}u, \quad m = 10, 20, \dots, 50,$$

by running the primal-dual algorithm of section 3.2. We have added the indices m to stress that the algorithm depends on it.

For the comparison with the three other reconstruction algorithms, we evaluate

$$e_{\text{wca}}^{(m)} = \max_{u \in \mathcal{M}_{\text{test}}} \|u - \widetilde{A}_{\text{wca}}^{(m)}(P_{W_m}u)\|, \quad m = 10, 20, \dots, 50.$$

We stress the fact that the three sets $\mathcal{M}_{\text{greedy}}$, $\widetilde{\mathcal{M}}$, and $\mathcal{M}_{\text{test}}$ are *different*. We compare this value with the performance of a straightforward reconstruction with the minimal V -norm recovery map,

$$e_{\text{mvn}}^{(m)} = \max_{u \in \mathcal{M}_{\text{test}}} \|u - P_{W_m}u\|, \quad m = 10, 20, \dots, 50,$$

with the mean square approach,

$$e_{\text{msa}}^{(m)} = \max_{u \in \mathcal{M}_{\text{test}}} \|u - \widetilde{A}_{\text{msa}}^{(m)}(P_{W_m}u)\|, \quad m = 10, 20, \dots, 50,$$

and with the best one-space affine algorithm,

$$e_{\text{one}}^{(m)} = \min_{1 \leq n \leq m} e_{\text{one}}^{(m,n)},$$

where

$$(3.28) \quad e_{\text{one}}^{(m,n)} = \max_{u \in \mathcal{M}_{\text{test}}} \|u - A_n^{(m)}(P_{W_m}u)\|, \quad m = 10, 20, \dots, 50.$$

Some remarks on the computation of the one-space algorithm are in order. First of all, we have used the average

$$\bar{u} := \frac{1}{\#\mathcal{M}_{\text{greedy}}} \sum_{u \in \mathcal{M}_{\text{greedy}}} u$$

as our offset. For $m \leq M$ and $n \leq m$ given, the one-space affine algorithm $A_n^{(m)}$ is the one involving the spaces W_m and $V_n = \bar{u} + V_n$, where $V_n = \text{span}\{u_1, \dots, u_n\}$. Its performance is given by $e_{\text{one}}^{(m,n)}$ in formula (3.28). Figure 3.5(a) shows $e_{\text{one}}^{(m,n)}$ as a function of n and m . Note that, for a fixed m , the error $e_{\text{one}}^{(m,n)}$ reaches a minimal value $e_{\text{one}}^m = e_{\text{one}}^{(m,n^*)}$ for a certain dimension $n^* = n^*(m)$ of the reduced model, given by a thick dot in the figure. This behavior is due to the trade-off between the increase of the approximation properties of $\widetilde{V}_n$ as n grows and the degradation of the stability of the algorithm, given by the increase of $\mu(\widetilde{V}_n, W_m)$ with n . For our comparison purposes, we use $A_{\text{one}}^{(m)} = A_{n^*(m)}^{(m)}$, that is, the best possible one-space algorithm based on the reduced basis spaces.

Figure 3.6 shows the reconstruction errors $e_{\text{wca}}^{(m)}$, $e_{\text{mvn}}^{(m)}$, $e_{\text{msa}}^{(m)}$, and $e_{\text{one}}^{(m)}$ of the four different approaches for $m = 10, 20, \dots, 50$. We also append a table with the values. We observe that a straightforward reconstruction with the minimal V -norm algorithm

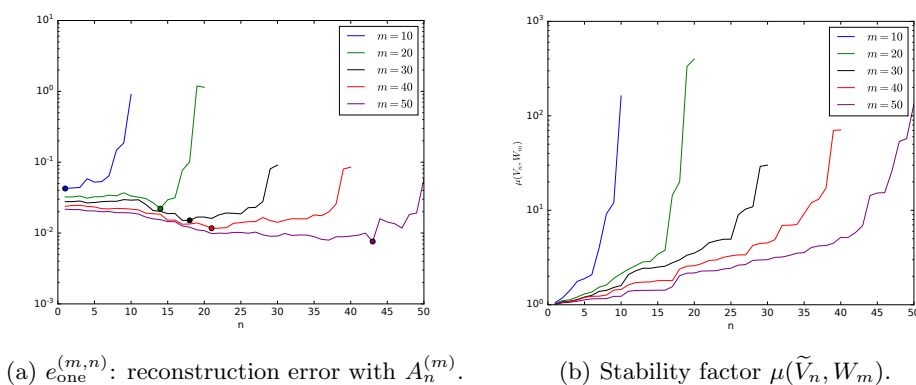
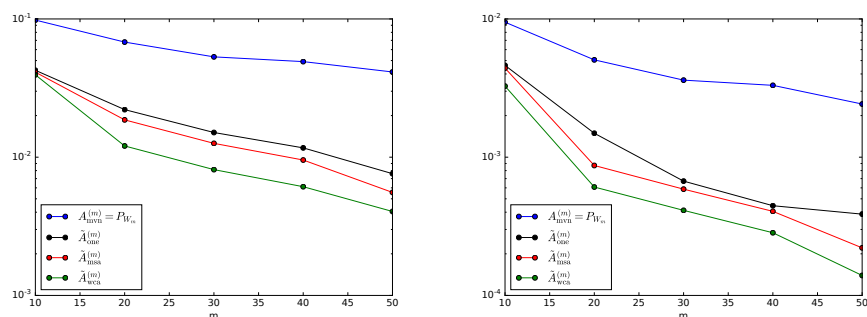


FIG. 3.5. One-space algorithm.

FIG. 3.6. Comparison of the reconstruction errors (left: $H_0^1(D)$ norm; right: $L^2(D)$ norm).

performs poorly in terms of approximation error and its quality improves only very mildly as we increase the number m of measurements. This justifies considering our three other, more sophisticated, reconstruction algorithms. In this respect, the results confirm first of all that $\tilde{A}_{\text{wca}}^{(m)}$ is the best reconstruction algorithm. The mean square approach appears to be slightly superior to the one-space algorithm but still worse than the best affine algorithm. Note that the accuracy improvement between the best affine algorithm and the one-space and mean square algorithms is of about a half-order of magnitude for each m .

Last but not least, we give some illustrations on the reconstruction algorithms applied to a particular snapshot function u from the test set $\mathcal{M}_{\text{test}}$. The target function is given in Figure 3.7 and Figures 3.8 and 3.9 show the resulting reconstructions of u from $P_{W_m} u$ with our four different algorithms and for $m = 20$ and 40. Visually, the reconstructed functions look very similar. However, the difference in quality can be better appreciated in the plots of the spatial errors $|u(\mathbf{x}) - A^{(m)}(u)(\mathbf{x})|$ as well as in the derivatives and their corresponding spatial errors.

Let us briefly discuss the complexity of the primal-dual algorithm. At each iteration of the algorithm, the main bottleneck is the computation of L^* (see (3.21)). It requires one to do J matrix-vector products with the matrices $Q_j \in \mathbb{R}^{N \times m(N+1)}$ and then do a summation of the resulting vectors. The cost of these operations thus increases linearly with J in terms of computational time and memory resources. In fact, the limitation in memory was the main reason to fix $J = 10^3$ and not work with a larger number of training snapshots. Let us make a quick count on the cost

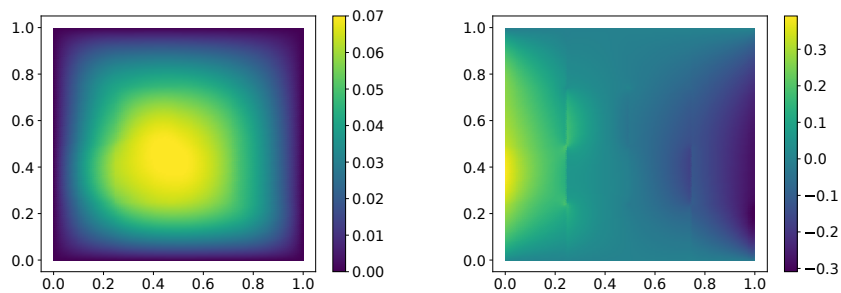
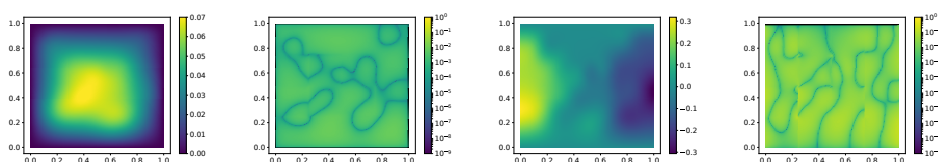
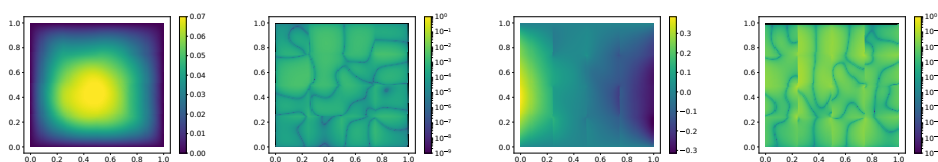


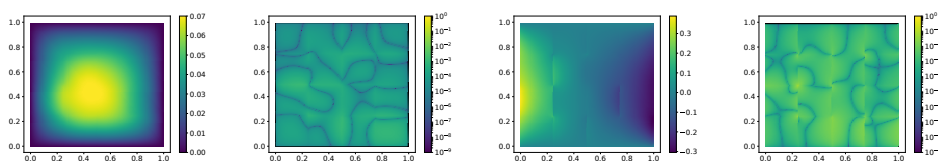
FIG. 3.7. Function u (left) and $\partial u/\partial x$ (right). The reconstruction of this function is given in Figures 3.8 and 3.9.



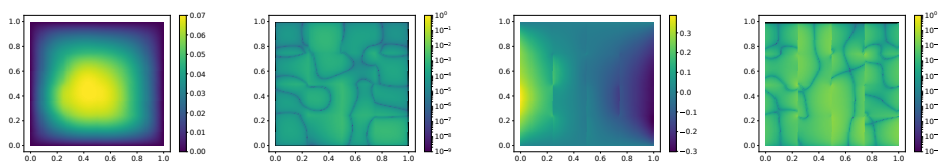
(a) Minimal V -norm: $P_{W_{20}}(u)$.



(b) One-space affine: $A_{\text{one}}^{(20)}(P_{W_{20}}(u))$



(c) Mean square algorithm: $\tilde{A}_{\text{msa}}^{(20)}(P_{W_{20}}(u))$



(d) Best affine: $\tilde{A}_{\text{wca}}^{(20)}(P_{W_{20}}(u))$

FIG. 3.8. Reconstruction of the given function u ($m = 20$). For each reconstruction strategy (i) the two first figures are $A^{(m)}(u)(\mathbf{x})$ and the spatial errors $|u(\mathbf{x}) - A^{(m)}(u)(\mathbf{x})|$; (ii) the two last figures are $\frac{\partial A^{(m)}(u)}{\partial x}(\mathbf{x})$ and the spatial errors $|\frac{\partial u}{\partial x}(\mathbf{x}) - \frac{\partial A^{(m)}(u)}{\partial x}(\mathbf{x})|$.

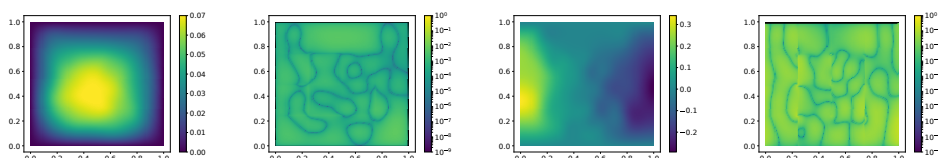
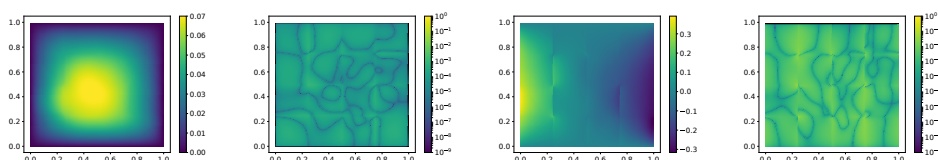
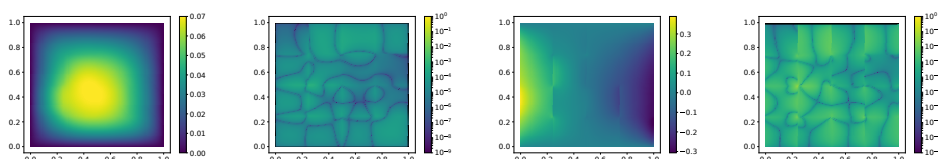
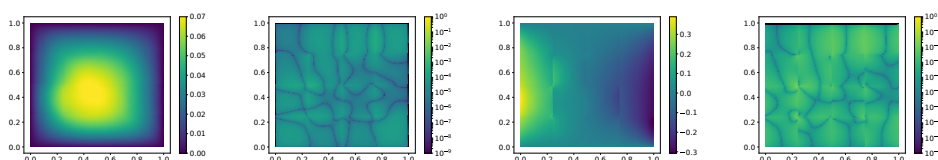
(a) Minimal V -norm: $P_{W_{40}}(u)$.(b) One-space affine: $A_{\text{one}}^{(40)}(P_{W_{40}}(u))$ (c) Mean square algorithm: $\tilde{A}_{\text{msa}}^{(40)}(P_{W_{40}}(u))$ (d) Best affine: $\tilde{A}_{\text{wca}}^{(40)}(P_{W_{40}}(u))$

FIG. 3.9. Reconstruction of the given function u ($m = 40$). For each reconstruction strategy (i) the two first figures are $A^{(m)}(u)(\mathbf{x})$ and the spatial errors $|u(\mathbf{x}) - A^{(m)}(u)(\mathbf{x})|$; (ii) the two last figures are $\frac{\partial A^{(m)}(u)}{\partial x}(\mathbf{x})$ and the spatial errors $|\frac{\partial u}{\partial x}(\mathbf{x}) - \frac{\partial A^{(m)}(u)}{\partial x}(\mathbf{x})|$.

in terms of the number of elements to store at each iteration. The matrices Q_j are sparse. For each row, there are $m + 1$ nonnegative coefficients. Therefore we need to store $N(m + 1)$ coefficients for each matrix; therefore, a total number of $JN(m + 1)$ coefficients. In our case, $N = 110$ was carefully fixed to guarantee that

$$\max_{u \in \mathcal{M}_{\text{greedy}} \cup \mathcal{M}_{\text{test}}} \|u - P_{V_N} u\| \leq \varepsilon_N = 5.10^{-5}.$$

We have m ranging between 10 and 50. Thus the number of nonnegative elements that we have to store for each Q_j ranges between 1210 and 5610. Therefore, taking $J = 10^3$ as in our computation, we need to handle a total number of coefficients ranging between $1.21 \cdot 10^6$ and $5.61 \cdot 10^6$.

4. Conclusions. In this paper, we have studied the notion of a best affine recovery map for a general state estimation problem, that is, the map A_{wca}^* that minimizes

the worst case error $E_{\text{wc}}(A, \mathcal{M})$ among all affine maps. This map is the solution to a convex optimization problem. Up to the additional perturbation induced from replacing $\mathcal{M}$ by a discrete training set $\tilde{\mathcal{M}}$, it can be efficiently computed by a primal-dual optimization algorithm. Since any affine recovery map is associated with a reduced basis V_n plus an offset $\bar{u}$, the optimal affine map amounts to applying the one-space method from [20] using an affine reduced model space $\bar{u} + V_n$ which is optimal for the reconstruction task. Our numerical tests confirm that this choice outperforms standard reduced basis spaces, which are not specifically constructed for the recovery problem, but rather for the approximation of $\mathcal{M}$.

Our approach is readily applicable to any type of parametric PDEs, ranging from linear PDEs with affine parameter dependence to nonlinear PDEs with nonaffine parameter dependence. We outline its main limitations:

- The first essential limitation lies in its confinement to linear or affine recovery algorithms. Let us stress that, while state estimation is a linear inverse problem in the sense that the observed data w are generated from u by a linear projection, the optimal recovery map is typically nonlinear, due to the complex nonlinear geometry of the solution manifold $\mathcal{M}$ that constitutes the prior. Therefore, going beyond the results provided by our method requires the development of nonlinear recovery strategies. One possible approach, currently under investigation, is to (i) consider a collection of affine reduced model spaces $\{\bar{u}_k + V^k : k = 1, \dots, K\}$, each of them of dimension $n_k \leq m$; (ii) use the observed data w to properly select a particular space from this collection; and (iii) apply the affine recovery algorithm using this particular data-dependent space. One standard way to obtain such local reduced model spaces is by splitting the parameter domain and searching for local reduced bases or proper orthogonal decomposition as, for example, proposed in [1] for forward modeling or in [15] for state estimation, as well as in [6]. However, the optimal affine recovery approach discussed in the present paper could also be used in order to improve on such constructions.
- The developed approach implicitly assumes that the parametric PDE model is perfect although, in full generality, the true physical state may not belong to $\mathcal{M}$. It is also assumed that measurements are noiseless. One way to readily extend this approach to the search of optimal affine maps that take into account model bias and measurement noise is as follows: suppose that the model bias is of size $\delta > 0$ in the sense that the real physical state u belongs to the offset

$$\mathcal{M}_\delta := \{v \in V : \exists y \in Y \text{ s.t. } \|v - u(y)\| \leq \delta\}.$$

Suppose further that measurements are given with some deterministic noise, that is, we are given $P_W u + \eta$ such that $\|\eta\| \leq \sigma$ for some noise level σ . Then, the optimal affine map is given by

$$\min_{A \text{ affine}} \max_{\substack{u \in \mathcal{M}_\delta, \\ \|\eta\| \leq \sigma}} \|u - A(P_W u + \eta)\|.$$

Once again we may emulate this optimization by introducing a discrete training set. Let us stress that such an optimization problem requires the knowledge of the size of the model bias δ and the noise level σ . While σ may be known for some applications, δ is very hard to estimate in practice.

An assessment of the obtained estimation accuracy relies, however, on the availability of computable bounds for the distance of the reduced spaces from the solution manifold which may depend on the type of the PDE model.

REFERENCES

- [1] D. AMSALLEM, C. FARHAT, AND M. J. ZAHR, *Nonlinear model order reduction based on local reduced order bases*, Internat. J. Numer. Methods Engrg., 92 (2012), pp. 891–916.
- [2] P. BINEV, A. COHEN, W. DAHMEN, R. DEVORE, G. PETROVA, AND P. WOJTASZCZYK, *Convergence rates for greedy algorithms in reduced basis methods*, SIAM J. Math. Anal., 43 (2011), pp. 1457–1472.
- [3] P. BINEV, A. COHEN, W. DAHMEN, R. DEVORE, G. PETROVA, AND P. WOJTASZCZYK, *Data assimilation in reduced modeling*, SIAM/ASA J. Uncertain. Quantif., 5 (2017), pp. 1–29.
- [4] P. BINEV, A. COHEN, O. MULA, AND J. NICHOLS, *Greedy algorithms for optimal measurements selection in state estimation using reduced models*, SIAM J. Uncertain. Quantif., 6 (2018), pp. 1101–1126.
- [5] B. BOJANOV, *Optimal recovery of functions and integrals*, in First European Congress of Mathematics, Vol. I, Paris, 1992, Progr. Math. 119, Birkhauser, Basel, 1994, pp. 371–390.
- [6] A. BONITO, A. COHEN, R. DEVORE, D. GUIGNARD, P. JANTSCH, AND G. PETROVA, *Nonlinear methods for model reduction*, ESAIM Math. Model. Numer. Anal., to appear, 2020, <https://doi.org/10.1051/m2an/2020057>.
- [7] A. BUFFA, Y. MADAY, A. T. PATERA, C. PRUD'HOMME, AND G. TURINICI, *A priori convergence of the greedy algorithm for the parameterized reduced basis*, ESAIM Math. Model. Numer. Anal., 46 (2012), pp. 595–603.
- [8] A. CHAMBOLE AND T. POCK, *A first-order primal-dual algorithm for convex problems with applications to imaging*, J. Math. Imaging Vision, 40 (2011), pp. 120–145.
- [9] A. COHEN, W. DAHMEN, R. DEVORE, AND J. NICHOLS, *Reduced Basis Greedy Selection Using Random Training Sets*, ESAIM Math. Model. Numer. Anal., 54 (2020), pp. 1509–1524.
- [10] A. COHEN AND R. DEVORE, *Approximation of high dimensional parametric PDEs*, Acta Numer., 24 (2015), pp. 1–159.
- [11] A. COHEN, R. DEVORE, AND C. SCHWAB, *Analytic regularity and polynomial approximation of parametric stochastic elliptic PDEs*, Anal. Appl., 9 (2011), pp. 11–47.
- [12] R. DEVORE, G. PETROVA, AND P. WOJTASZCZYK, *Greedy algorithms for reduced bases in Banach spaces*, Constr. Approx., 37 (2013), pp. 455–466.
- [13] M. DASHTI AND A. M. STUART, *The Bayesian approach to inverse problems*, Handbook of Uncertainty Quantification, R. Ghanem, D. Higdon, and H. Owhadi, eds., Springer, Cham, Switzerland, 2015.
- [14] J. FADILI AND O. MULA, *Primal-Dual Splitting for the Max of Convex Functions*, in progress.
- [15] F. GALARCE, J.-F. GERBEAU, D. LOMBARDI, AND O. MULA, *State Estimation with Nonlinear Reduced Models Application to the Reconstruction of Blood Flows with Doppler Ultrasound Images*, preprint, arXiv:1904.13367, 2019.
- [16] P. L. HOUTEKAMER AND H. L. MITCHELL, *Ensemble Kalman filtering*, Quart. J. R. Meteorol. Soc., 131 (2005), pp. 3269–3289.
- [17] M. KÄRCHER, S. BOYAVAL, M. A. GREPL, AND K. VEROY, *Reduced basis approximation and a posteriori error bounds for 4D-Var data assimilation*, J. Optim. Engrg., 19 (2018), pp. 663–695.
- [18] K. LAW, A. STUART, AND K. ZYGALAKIS, *Data Assimilation, A Mathematical Introduction*, Texts in Appl. Math. 62, Springer, Cham, Switzerland, 2015.
- [19] A. C. LORENC, *A global three-dimensional multivariate statistical interpolation scheme*, Mon. Weather Rev., 109 (1981), pp. 701–721.
- [20] Y. MADAY, A. T. PATERA, J. D. PENN, AND M. YANO, *A parametrized-background data-weak approach to variational data assimilation: Formulation, analysis, and application to acoustics*, Internat. J. Numer. Methods Engrg., 102 (2015), pp. 933–965.
- [21] C. A. MICCHELLI AND T. J. RIVLIN, *Lectures on optimal recovery*, in Numerical analysis, Lancaster, England, 1984, Lecture Notes in Math. 1129, Springer, Berlin, 1985, pp. 21–93.
- [22] E. NOVAK AND H. WOZNIAKOWSKI, *Tractability of Multivariate Problems, Volume I: Linear Information*, EMS Tracts Math. 6, European Mathematical Society, Zurich, 2008.
- [23] G. ROZZA, D. B. P. HUYNH, AND A. T. PATERA, *Reduced basis approximation and a posteriori error estimation for affinely parametrized elliptic coercive partial differential equations $\tilde{N}$ application to transport and continuum mechanics*, Arch. Comput. Methods Engrg., 15 (2008), pp. 229–275.

- [24] S. SEN, *Reduced-basis approximation and a posteriori error estimation for many-parameter heat conduction problems*, Numer. Heat Transf. B, 54 (2008), pp. 369–389.
- [25] A. M. STUART, *Inverse problems: A Bayesian perspective*, Acta Numer., 19 (2010), pp. 451–559.
- [26] T. TADDEI, *An adaptive parametrized-background data-weak approach to variational data assimilation*, ESAIM Math. Model. Numer. Anal., 51 (2017), pp. 1827–1858.
- [27] B. PEHERSTORFER, K. WILLCOX, AND M. GUNZBURGER, *Survey of multifidelity methods in uncertainty propagation, inference, and optimization*, SIAM Rev., 60 (2018), pp. 550–591.