



An instance optimality property for approximation problems with multiple approximation subspaces

Cedric Herzet, Mamadou Diallo, Patrick Héas

► To cite this version:

Cedric Herzet, Mamadou Diallo, Patrick Héas. An instance optimality property for approximation problems with multiple approximation subspaces. 2018. hal-01913339

HAL Id: hal-01913339

<https://inria.hal.science/hal-01913339>

Preprint submitted on 6 Nov 2018

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

An instance optimality property for approximation problems with multiple approximation subspaces

C. Herzet^{1,2} · M. Diallo¹ · P. Héas¹

Received: date / Accepted: date

Abstract Model-order reduction methods tackle the following general approximation problem: find an “easily-computable” but accurate approximation $\hat{\mathbf{h}}$ of some target solution $\mathbf{h}^*$. In order to achieve this goal, standard methodologies combine two main ingredients: *i*) a set of problem-specific constraints; *ii*) some “simple” prior model on the set of target solutions. The most common prior model encountered in the literature assume that the target solution $\mathbf{h}^*$ is “close” to some low-dimensional subspace. Recently, triggered by the work by Binev *et al.* [5], several contributions have shown that refined prior models (based on a *set* of embedded approximation subspaces) may lead to enhanced approximation performance. Unfortunately, to date, no theoretical results have been derived to support the good empirical performance observed in these contributions. The goal of this work is to fill this gap. More specifically, we provide a mathematical characterization of the approximation performance achievable by some particular “multi-space” decoder and emphasize that, in some specific setups, this “multi-space” decoder has provably better recovery guarantees than its standard counterpart based on a single approximation subspace.

Keywords Model-order reduction · Petrov-Galerkin approximation · Multi-space approximation · Instance optimal property

The authors thank the “Agence nationale de la recherche” for its financial support through the Geronimo project (ANR-13-JS03-0002).

C. Herzet
Tel.: +33-2-99847350
E-mail: cedric.herzet@inria.fr

¹ INRIA, Centre Rennes Bretagne-Atlantique, Campus de Beaulieu, 35000 Rennes, France

² IMT Atlantique, Lab-STICC, UBL, 29238 Brest, France

1 Introduction

Many approximation methods encountered in the domain of model-order reduction are based on the following abstract problem:

$$\forall \mathbf{h}^* \in \mathcal{M}, \text{ find a "proper" } \hat{\mathbf{h}} \in \mathcal{M}_{\text{prior}} \cap \mathcal{P}_{\mathbf{h}^*} \quad (1)$$

where $\mathcal{M}$ is a set of target solutions, $\mathcal{M}_{\text{prior}}$ corresponds to some prior knowledge about $\mathcal{M}$, and $\mathcal{P}_{\mathbf{h}^*}$ is a set of elements satisfying some problem-specific linear constraints:

$$\mathcal{P}_{\mathbf{h}^*} = \{\mathbf{h} : \langle \mathbf{h}', \mathbf{h} \rangle = \langle \mathbf{h}', \mathbf{h}^* \rangle \forall \mathbf{h}' \in W_m\} \quad (2)$$

where W_m is some m -dimensional linear subspace of the ambient Hilbert space $\mathcal{H}$.¹ For example, “projection-based” reduction of linear parametrized partial differential equations (PPDE) [11], some approximation methods for non-linear operators (*e.g.*, EIM [3], DEIM [6], Gappy POD [7]) or data-assimilation schemes with reduced-order models [5] can be recast as particular instances of (1).

In these problems, the choice of $\mathcal{M}_{\text{prior}}$ and $\mathcal{P}_{\mathbf{h}^*}$ plays a crucial role since it determines the trade-off between accuracy and complexity achievable by the approximation method. A standard option consists of choosing $\mathcal{M}_{\text{prior}}$ as:

$$\mathcal{M}_{\text{prior}} = \{\mathbf{h} : \text{dist}(\mathbf{h}, V_n) \leq \hat{\epsilon}_n\} \quad (3)$$

where V_n is a n -dimensional subspace, $\text{dist}(\mathbf{h}, V_n) \triangleq \min_{\mathbf{h}' \in V_n} \|\mathbf{h} - \mathbf{h}'\|$ and $\hat{\epsilon}_n \geq 0$. A variety of methods have been proposed in the literature to select “good” subspaces W_m and V_n , see *e.g.*, [4, 11]. These methodologies are grounded on the following “instance optimality property” (or some variation thereof) which is valid for most of the decoders (1) exploiting constraints of the form (2)-(3):

$$\|\hat{\mathbf{h}} - \mathbf{h}^*\| \leq C(W_m, V_n) \text{dist}(\mathbf{h}^*, V_n) \quad (4)$$

where $C(W_m, V_n)$ is some constant depending on W_m and V_n .

Recently, triggered by the work by Binev *et al.* [5], more refined definitions of $\mathcal{M}_{\text{prior}}$ have been considered in [1, 8–10]. In these works, $\mathcal{M}_{\text{prior}}$ is built from a sequence of embedded approximation subspaces, that is

$$\mathcal{M}_{\text{prior}} = \cap_{k=0}^n \{\mathbf{h} : \text{dist}(\mathbf{h}, V_k) \leq \hat{\epsilon}_k\} \quad (5)$$

where $\hat{\epsilon}_k \geq 0$, $\dim(V_k) = k$ and

$$V_0 \subset V_1 \subset \dots \subset V_n. \quad (6)$$

In [5], the authors select the subspaces $\{V_k\}_{k=0}^n$ and widths $\{\hat{\epsilon}_k\}_{k=0}^n$ so that $\mathcal{M} \subseteq \mathcal{M}_{\text{prior}}$, and propose an iterative procedure to identify a point of $\mathcal{M}_{\text{prior}} \cap$

¹ We assume that all the quantities defined above belong to a Hilbert space $\mathcal{H}$ with inner product $\langle \cdot, \cdot \rangle$ and induced norm $\|\cdot\|$.

$\mathcal{P}_{\mathbf{h}^*}$. The authors also derive an “instance optimal property” bounding the approximation error as a function of the distance between $\mathbf{h}^*$ and the subspaces $\{V_k\}_{k=0}^n$. From a practical point of view, the algorithm proposed in [5] has a complexity scaling as $\mathcal{O}((m+n)^2)$ (per iteration) but requires the knowledge of some orthonormal bases of W_m and $W_m \oplus V_n$. When the subspace W_m is fixed, the identification of these bases can be done once for all in advance, with no impact on the online complexity of the procedure. In the more general case where the definition of W_m may vary with $\mathbf{h}^* \in \mathcal{M}$ (as it may occur in the case of model-order reduction of PPDE’s), the identification of these bases may however become computationally too expensive.

In order to circumvent this problem, a slightly different definition of $\mathcal{M}_{\text{prior}}$ have been considered in several recent works, see [1, 8–10]. In these contributions, the authors relax the constraint “ $\hat{\mathbf{h}} \in \mathcal{P}_{\mathbf{h}^*}$ ” but impose the approximation $\hat{\mathbf{h}}$ to belong to the n -dimensional subspace V_n .² As pointed out in [9, 10], imposing “ $\hat{\mathbf{h}} \in V_n$ ” allows to implement approximation algorithms whose complexity scales as $\mathcal{O}((m+n)^2)$ per iteration without requiring the evaluation of orthonormal bases depending on W_m . This approach thus allows for slightly more computational flexibility than the procedure proposed in [5]. As shown in these works, this particular definition for $\mathcal{M}_{\text{prior}}$ leads to great improvements of the approximation performance (as compared to that induced by the standard prior (3)) in both the fields of model-order reduction of PPDEs and the approximation of non-linear operators.

In this paper, we focus on the approximation procedure considered in [1, 8–10]. Although numerical assessments have shown the relevance of this approach, to date, no theoretical guarantees have been made available to support these empirical evidences. The goal of this work is to fill this gap. We derive an instance optimal property relating the approximation performance of the decoder considered in [1, 8–10] to the distances between the target solution $\mathbf{h}^*$ and the approximation subspaces $\{V_k\}_{k=0}^n$. We show on several examples that, for particular choices of subspaces $\{V_k\}_{k=0}^n$ and W_m , this enhanced prior can lead to much better approximation guarantees than its standard counterpart (3).

The paper is organized as follows. In Section 2, we recall the standard methodologies encountered in the literature when $\mathcal{M}_{\text{prior}}$ is defined as in (3). In Section 3, we give a precise definition of the “multi-space” decoder considered in [1, 8–10] and state an optimal instance property characterizing its performance. We particularize our result to two different setups and show that the “multi-space” decoder has more favorable guarantees of performance than the “single-space” approach (3) in these cases. The proof of our main result is finally detailed in Section 4.

² We give a precise formulation of the approximation problem considered in [1, 8–10] in Section 3 of this paper.

2 The Single-space Approximation

In this section, we consider a particular concrete example of the abstract problem stated in (1) when $\mathcal{M}_{\text{prior}}$ is defined as in (3). Since this prior model only involves one approximation subspace V_n , we will refer to it as “single-space approximation” in the sequel. The material presented in this section is not novel but is intended to support our discussion in the next section.

When $\mathcal{M}_{\text{prior}}$ is defined as in (3) and $m = n$, a standard approach to select an element of $\mathcal{M}_{\text{prior}} \cap \mathcal{P}_{\mathbf{h}^*}$ is as follows:

$$\text{Find } \hat{\mathbf{h}}_{\text{SS}} \in V_n \text{ such that } \langle \mathbf{w}_j, \hat{\mathbf{h}}_{\text{SS}} \rangle = y_j \text{ for all } j = 1 \dots n, \quad (7)$$

where $\{\mathbf{w}_j\}_{j=1}^m$ is a basis of W_m and $y_j \triangleq \langle \mathbf{w}_j, \mathbf{h}^* \rangle$. This problem can also be expressed³ in a variational form as

$$\text{Find } \hat{\mathbf{h}}_{\text{SS}} \in \arg \min_{\mathbf{h} \in V_n} \sum_{j=1}^n (y_j - \langle \mathbf{w}_j, \mathbf{h} \rangle)^2. \quad (8)$$

As shown in [5, Section 2], the approximation computed in (7) can in fact be seen as a “min-max” solution of (1): under some mild non-degeneracy conditions, (7) can indeed be rewritten as

$$\hat{\mathbf{h}}_{\text{SS}} = \arg \min_{\mathbf{h} \in \mathcal{M}_{\text{prior}} \cap \mathcal{P}_{\mathbf{h}^*}} \sup_{\mathbf{h}' \in \mathcal{M}_{\text{prior}} \cap \mathcal{P}_{\mathbf{h}^*}} \|\mathbf{h} - \mathbf{h}'\|. \quad (9)$$

Hence, as long as prior (3) is considered, $\hat{\mathbf{h}}_{\text{SS}}$ corresponds to the element of $\mathcal{M}_{\text{prior}} \cap \mathcal{P}_{\mathbf{h}^*}$ minimizing the worst-case approximation error.

Problem (7) is the basis of many methodologies encountered in the field of model-order reduction. In the context of “projection-based” reduction of PPDEs, (7) is better known as “Petrov-Galerkin projection”. In this case, $\mathcal{M}$ corresponds to a set of solutions of a differential equation, *i.e.*,

$$\mathcal{M} = \{\mathbf{h}^* : \text{PDE}(\mathbf{h}^*, \theta) = 0 \text{ for some } \theta \in \Theta\} \quad (10)$$

where Θ is some set of parameters and $\text{PDE}(\mathbf{h}^*, \theta) = 0$ is an abstract notation for the PPDE; V_n is an approximation subspace for $\mathcal{M}$ and W_m derives from the form defining the weak formulation of the PPDE, see *e.g.*, [11].

Many procedures to approximate the output of a non-linear operator (*e.g.*, EIM [3], DEIM [6], Gappy POD [7]) may also be seen as particular instances of (7). Here, $\mathcal{M}$ corresponds to a set of outputs of an operator $L : \Theta \rightarrow \mathcal{H}$, *i.e.*,

$$\mathcal{M} = \{\mathbf{h}^* = L(\theta) : \theta \in \Theta\}, \quad (11)$$

V_n is an approximation subspace for $\mathcal{M}$ and W_m is commonly chosen so that $y_j = \langle \mathbf{w}_j, L(\theta) \rangle$ can be computed efficiently $\forall \theta \in \Theta$.

³ Equivalence holds as long as a solution to (7) exists.

Finally, in [5], the authors considered a problem related to (7) in the domain of data assimilation with reduced-order models. In this context, $\mathcal{M}$ typically represent the set of solutions of a PPDE as in (10), V_n is an approximation subspace for $\mathcal{M}$ and W_m defines the observation operator used in the data assimilation process.

The performance of the single space approximation (7) is characterized by the well-known Babuska's theorem [12] stated below. The statement of this theorem involves the singular values of the following Gram matrix

$$\mathbf{G} = [\langle \mathbf{w}_i, \mathbf{v}_j \rangle]_{i,j} \in \mathbb{R}^{n \times n}. \quad (12)$$

where $\{\mathbf{v}_j\}_{j=1}^n$ is an arbitrary basis of V_n .

Theorem 1 (Babuska's Theorem) *Let σ_1 and σ_n respectively denote the largest and smallest singular values of the Gram matrix defined in (12). If $\sigma_n > 0$ then (7) has a unique solution which satisfies*

$$\|\mathbf{h}^* - \hat{\mathbf{h}}_{\text{SS}}\| \leq \frac{\sigma_1}{\sigma_n} \text{dist}(\mathbf{h}^*, V_n). \quad (13)$$

We note that (13) is only one particular example of instance optimal property valid for single-space approximations in Hilbert spaces. Other versions of instance optimal properties exist. For example in [5, Theorem 2.9], a more general formulation of (13) has been derived in Hilbert spaces when $m > n$. In Banach spaces, (4) holds with a slightly larger constant $C(W_n, V_n) = 1 + \frac{\sigma_1}{\sigma_n}$, see [2]. In the field of non-linear operator approximation, instance optimal properties involving L_∞ -norms have been derived, see *e.g.*, [3]. In what follows, we will restrict our attention to (13), since it is valid in the context of Hilbert spaces considered in this paper, and is more amenable to comparisons with our main result stated in Theorem 2.

3 The Multi-space Approximation

In this section, we present a theoretical result supporting the performance of the “multi-space” decoder considered in [1, 8–10]. We state our main result in Theorem 2 and provide two examples of scenarios in which the multi-space decoder have provably better performance guarantees than the standard “single space” decoder (7).

Before stating our result, we recall the definition of the multi-space decoder considered in [1, 8–10]:⁴

$$\begin{aligned} \text{Find } \hat{\mathbf{h}}_{\text{MS}} \in \arg \min_{\mathbf{h} \in V_n} \sum_{j=1}^m (y_j - \langle \mathbf{w}_j, \mathbf{h} \rangle)^2 \\ \text{subject to } \text{dist}(\mathbf{h}, V_k) \leq \hat{\epsilon}_k, \quad k = 0 \dots n. \end{aligned} \quad (14)$$

⁴ In this paper we assume that the constraints are defined $\forall k \in \{0 \dots n\}$. All the derivations presented in this paper may nevertheless be easily extended to the case where the constraints in (14) are only available for *some* $k \in \{0 \dots n\}$.

In the sequel, we assume without loss of generality that the vectors $\{\mathbf{w}_j\}_{j=1}^m$ are linearly independent. We also suppose that the subspaces $\{V_k\}_{k=0}^n$ are embedded, that is, obey (6).

We note that (14) can be seen as an extension of the standard Petrov-Galerkin approach discussed in Section 2. More specifically, if one removes the constraints in (14) and let $m = n$, the problem becomes

$$\text{Find } \hat{\mathbf{h}} \in \arg \min_{\mathbf{h} \in V_n} \sum_{j=1}^n (y_j - \langle \mathbf{w}_j, \mathbf{h} \rangle)^2, \quad (15)$$

i.e., is equivalent to problem (8). On the other hand, we note that the constraints in (14) define a set of feasible points of the form (5). Hence, if the subspaces $\{V_k\}_{k=0}^n$ and widths $\{\hat{\epsilon}_k\}_{k=0}^n$ are properly chosen, these constraints add some valuable information about the position of the sought solution $\mathbf{h}^*$ in $\mathcal{H}$. We may thus expect the multi-space decoder to lead to enhanced performance in some specific situations. In this section, we provide a mathematical support to this intuition.

In the following, we will make the assumption that the constraints in (14) are such that

$$\forall \mathbf{h}^* \in \mathcal{M} : \text{dist}(\mathbf{h}^*, V_k) \leq \hat{\epsilon}_k \text{ for all } k = 0 \dots n, \quad (16)$$

that is, they are satisfied by any target solution $\mathbf{h}^* \in \mathcal{M}$. Under this assumption, we provide a mathematical characterization of the performance achievable by the multi-space decoder (14). More specifically, we derive an instance optimality property relating the approximation error $\|\hat{\mathbf{h}}_{\text{MS}} - \mathbf{h}^*\|$ to the distance between $\mathbf{h}^*$ and the different approximation subspaces $\{V_k\}_{k=0}^n$. Our result is presented in Theorem 2 below.

In order to state our result we need to introduce the following quantities. We first define the short-hand notations

$$\epsilon_k = \text{dist}(\mathbf{h}^*, V_k) \text{ for all } k = 0 \dots n, \quad (17)$$

and

$$\gamma = \sup_{\mathbf{h} \in V_n^\perp, \|\mathbf{h}\|=1} \left(\sum_{j=1}^m \langle \mathbf{w}_j, \mathbf{h} \rangle^2 \right)^{\frac{1}{2}}, \quad (18)$$

where $V_n^\perp$ is the orthogonal complement of V_n in $\mathcal{H}$. We let $\{\mathbf{v}_j\}_{j=1}^n$ be an orthonormal basis of V_n such that

$$V_k = \text{span} \left(\{\mathbf{v}_j\}_{j=1}^k \right). \quad (19)$$

We note that such a basis always exists since we assume that the sequence of subspaces $\{V_k\}_{k=0}^n$ obeys (6).

We define the Gram matrix $\mathbf{G}$ as in (12) and let

$$\delta_j = \sum_{k=1}^n |x_{kj}|(\hat{\epsilon}_{k-1} + \epsilon_{k-1}), \quad (20)$$

where x_{kj} are the elements of matrix $\mathbf{X}$ appearing in the singular value decomposition of $\mathbf{G}$, that is $\mathbf{G} = \mathbf{U}\mathbf{A}\mathbf{X}^T$, where $\mathbf{U} \in \mathbb{R}^{m \times m}$, $\mathbf{X} \in \mathbb{R}^{n \times n}$ are orthogonal matrices and $\mathbf{A} \in \mathbb{R}^{m \times n}$ is the diagonal matrix of singular values $\{\sigma_j\}_{j=1}^{\min(m,n)}$. In the sequel we will consider the extended set $\{\sigma_j\}_{j=1}^n$ by using the following convention: if $n > m$, we define $\sigma_j = 0$ for all $j > m$. Without loss of generality, we assume that the singular values $\{\sigma_j\}_{j=1}^n$ are sorted by decreasing order of magnitude.

Using these notations, our result reads:

Theorem 2 *Assume $\mathbf{h}^*$ verifies (16) and let $y_j = \langle \mathbf{w}_j, \mathbf{h}^* \rangle$ for $j = 1 \dots m$. Then any solution $\hat{\mathbf{h}}_{\text{MS}}$ of (14) verifies*

$$\|\mathbf{h}^* - \hat{\mathbf{h}}_{\text{MS}}\| \leq \begin{cases} \left(\sum_{j=\ell+1}^n \delta_j^2 + \rho \delta_\ell^2 + \epsilon_n^2 \right)^{\frac{1}{2}} & \text{if } \sum_{j=1}^n \sigma_j^2 \delta_j^2 \geq 4\gamma^2 \epsilon_n^2, \\ \left(\sum_{j=1}^n \delta_j^2 + \epsilon_n^2 \right)^{\frac{1}{2}} & \text{otherwise,} \end{cases} \quad (21)$$

where ℓ is the largest integer such that

$$\sum_{j=\ell}^n \sigma_j^2 \delta_j^2 \geq 4\gamma^2 \epsilon_n^2, \quad (22)$$

and $\rho \in [0, 1]$ is defined as

$$\rho \sigma_\ell^2 \delta_\ell^2 + \sum_{j=\ell+1}^n \sigma_j^2 \delta_j^2 = 4\gamma^2 \epsilon_n^2. \quad (23)$$

Moreover, if $\sigma_n > 0$, (14) admits a unique solution.

The proof of Theorem 2 is reported in Section 4. We note that the structure of the instance optimal property stated in Theorem 2 is similar to that derived by Binev *et al.* in [5, Theorem 3.2] (for a different multi-space decoder) although it involves different constants and the singular values of a different Gram matrix. Scrutinizing the proof of our result in Section 4, we notice nevertheless that the reasoning leading to Theorem 2 is quite different from that developed in [5]. This is due to the fact that the formulations of the decoders considered here and in [5] are quite different and thus necessitate distinct developments to derive the corresponding instance optimal property.

Because they involve singular values of different Gram matrices, the recovery results in [5, Theorem 3.2] and Theorem 2 can only be compared when $\{\mathbf{w}_j\}_{j=1}^m$ is an orthonormal basis. In that particular case, it can be seen that the result in [5, Theorem 3.2] is slightly more favorable (to some constant factor)

than that in Theorem 2.⁵ This observation could have been intuitively expected since the prior information exploited in (14) correspond to a degraded version of that used in [5]. Indeed, as mentioned previously, the multi-space decoder in [5] is a particular instance of (1) with $\mathcal{P}_{\mathbf{h}^*}$ and $\mathcal{M}_{\text{prior}}$ defined in (2) and (5) respectively; on the other hand, the multi-space decoder (14), although also considering a prior of the form (5), imposes $\hat{\mathbf{h}}_{\text{MS}} \in V_n$ and relaxes the constraint $\hat{\mathbf{h}}_{\text{MS}} \in \mathcal{P}_{\mathbf{h}^*}$. These modifications remove some relevant information about the position of $\mathbf{h}^*$ in $\mathcal{H}$ (since $\mathbf{h}^* \in \mathcal{P}_{\mathbf{h}^*}$ by construction and $\mathbf{h}^*$ does not necessarily belong to V_n). We note however that, as mentioned in the introduction of this paper, these modifications also allow for a more flexible computational implementation of the decoder in some situations.

We now turn our attention to the comparison between the performance achievable by the standard single-space decoder (7) and its multi-space version (14). More specifically, we show that the recovery guarantees obtained in the multi-space setup may be much more favorable than in the single-space setup in some specific settings. In order to allow a comparison between the single and multi-space decoders, we consider the case where $m = n$ and assume that $\{\mathbf{w}_j\}_{j=1}^m$ is an orthonormal basis. We note that, in such a case, we have $\sigma_1 \leq 1$ and $\gamma \leq 1$. The specific settings considered hereafter are inspired from [5] and are described in Examples 1 and 2. They correspond to different choices of matrix $\mathbf{X}$ appearing in the singular value decomposition of the Gram matrix $\mathbf{G}$ defined in (12). Since the latter matrix directly depends on the bases $\{\mathbf{w}_j\}_{j=1}^m$ and $\{\mathbf{v}_j\}_{j=1}^n$, these examples thus correspond to some particular choices of the observation and approximation subspaces W_m and $\{V_k\}_{k=0}^n$.

Example 1 We first assume that $\mathbf{X} = \mathbf{I}_n$ in the singular-value decomposition of $\mathbf{G}$. We set $\hat{\epsilon}_j = \epsilon_j$ and assume that

$$\epsilon_j = \begin{cases} 1 & j = 0 \dots n-3, \\ \epsilon^{\frac{1}{2}} & j = n-2, n-1, \\ \epsilon & j = n, \end{cases} \quad (24)$$

for some $\epsilon \ll 1$. Moreover, we let

$$\sigma_j = \begin{cases} 1 & j = 1 \dots n-3, \\ \epsilon^{\frac{1}{2}} & j = n-2, n-1, \\ \epsilon & j = n. \end{cases} \quad (25)$$

In this setup, the upper bound (13) of Theorem 1 becomes:

$$\|\hat{\mathbf{h}}_{\text{SS}} - \mathbf{h}^*\| \leq \sigma_n^{-1} \text{dist}(\mathbf{h}^*, V_n) = \epsilon^{-1} \epsilon = 1. \quad (26)$$

On the other hand, because $\mathbf{X} = \mathbf{I}_n$ and $\hat{\epsilon}_j = \epsilon_j$, we have

$$\delta_j = \hat{\epsilon}_{j-1} + \epsilon_{j-1} = 2\epsilon_{j-1}. \quad (27)$$

⁵ More specifically, the factor $4\gamma^2$ in Theorem 2 is equal to 1 in [5, Theorem 3.2].

The index ℓ appearing in Theorem 2 is smaller or equal to $n - 1$ since

$$\begin{aligned}\sigma_n^2 \delta_n^2 &= \sigma_n^2 (2\epsilon_{n-1})^2 = 4\epsilon^3 \ll 4\epsilon^2, \\ \sigma_{n-1}^2 \delta_{n-1}^2 &= \sigma_{n-1}^2 (2\epsilon_{n-2})^2 = 4\epsilon^2,\end{aligned}$$

and thus

$$\sigma_{n-1}^2 \delta_{n-1}^2 + \sigma_n^2 \delta_n^2 \geq 4\epsilon^2 \geq 4\gamma^2 \epsilon^2 \quad (28)$$

since $\gamma \leq 1$. The upper bound in Theorem 2 becomes

$$\begin{aligned}\|\mathbf{h}^* - \hat{\mathbf{h}}_{\text{MS}}\| &\leq (\delta_{n-1}^2 + \delta_n^2 + \epsilon_n^2)^{\frac{1}{2}}, \\ &= (4\epsilon + 4\epsilon + \epsilon^2)^{\frac{1}{2}}, \\ &\leq 3\epsilon^{\frac{1}{2}}.\end{aligned} \quad (29)$$

Hence the bound in the multi-space setup (29) can be arbitrarily small as compared to (26) when $\epsilon \rightarrow 0$. $\square$

Example 2 We now consider $\mathbf{X} = n^{-\frac{1}{2}} \mathbf{1}_{n \times n}$ where $\mathbf{1}_{n \times n}$ is an $n \times n$ matrix of 1's. We set $\hat{\epsilon}_j = \epsilon_j$ and assume that

$$\epsilon_j = \begin{cases} \frac{1}{2} & j = 0, \\ \frac{1}{2(n-1)} & j = 1 \dots n-1, \\ \epsilon & j = n, \end{cases} \quad (30)$$

for some $\epsilon \ll n^{-1}$ (Note that we must have: $\epsilon \leq \frac{1}{2(n-1)}$ by definition). Moreover, we let

$$\sigma_j = \begin{cases} \sigma & j = 1 \dots n-1, \\ \epsilon^2 & j = n, \end{cases} \quad (31)$$

for some $1 \geq \sigma > \epsilon$ whose value will be specified below.

With these choices, the upper bound (13) of Theorem 1 becomes:

$$\|\hat{\mathbf{h}}_{\text{SS}} - \mathbf{h}^*\| \leq \sigma_n^{-1} \text{dist}(\mathbf{h}^*, V_n) = \epsilon^{-2} \epsilon = \epsilon^{-1}. \quad (32)$$

On the other hand, we have

$$\begin{aligned}\delta_j &= \sum_{k=1}^n |x_{kj}| (\hat{\epsilon}_{k-1} + \epsilon_{k-1}), \\ &= 2n^{-\frac{1}{2}} \sum_{k=1}^n \epsilon_{k-1}, \\ &= 2n^{-\frac{1}{2}}.\end{aligned} \quad (33)$$

By choosing σ such that (we remind the reader that $\sigma_{n-1} = \sigma$ by definition (31))

$$\sigma_{n-1}^2 \delta_{n-1}^2 + \sigma_n^2 \delta_n^2 = 4\epsilon^2, \quad (34)$$

we obtain that index ℓ appearing in Theorem 2 is smaller or equal to $n - 1$ since $\gamma \leq 1$. The upper bound in Theorem 2 then reads

$$\begin{aligned} \|\mathbf{h}^* - \hat{\mathbf{h}}_{\text{MS}}\| &\leq (\delta_{n-1}^2 + \delta_n^2 + \epsilon_n^2)^{\frac{1}{2}}, \\ &= (4n^{-1} + 4n^{-1} + \epsilon^2)^{\frac{1}{2}}, \\ &\leq 3n^{-\frac{1}{2}}, \end{aligned} \quad (35)$$

where the last inequality follows from our initial assumption $\epsilon \ll n^{-1}$. $\square$

We conclude this section by providing a graphical representation of the geometry of problem (14). Fig. 1 gives an illustration of the feasible set and the iso-contours of the cost function appearing in (14); these quantities are plotted in the plane V_n for $m = n = 2$.

If $\sigma_n > 0$, it can be seen (see Appendix A) that the cost function $f(\mathbf{h}) \triangleq \sum_{j=1}^n (y_j - \langle \mathbf{w}_j, \mathbf{h} \rangle)^2$ can also be rewritten for any $\mathbf{h} \in V_n$ as

$$f(\mathbf{h}) = \sum_{j=1}^n \sigma_j^2 \left(\langle \mathbf{v}_j^*, \hat{\mathbf{h}}_{\text{SS}} \rangle - \langle \mathbf{v}_j^*, \mathbf{h} \rangle \right)^2, \quad (36)$$

where

$$\mathbf{v}_j^* = \sum_{i=1}^n x_{ij} \mathbf{v}_i. \quad (37)$$

Here, the elements x_{ij} 's correspond to the components of the orthonormal matrix $\mathbf{X}$ appearing in the singular value decomposition of $\mathbf{G}$. From (36), we thus see that the iso-contours of $f(\mathbf{h})$ in V_n correspond to n -dimensional ellipsoids with center equal to $\hat{\mathbf{h}}_{\text{SS}}$ and principal axes equal to $\{\mathbf{v}_j^*\}_{j=1}^n$. Moreover, the elongation of the ellipsoids along each axis $\mathbf{v}_j^*$ is inversely proportional to the singular value σ_j .

From a geometric point of view, bad recovery guarantees in the single-space setup (*i.e.*, large value for $\frac{\sigma_1}{\sigma_n}$) corresponds to situations where the “iso-contour” ellipsoids are much more elongated along (at least) one direction than another: the center of the ellipsoid $\hat{\mathbf{h}}_{\text{SS}}$ may then be quite distant from the optimal orthogonal projection $P_{V_n}(\mathbf{h}^*)$.

The prior information used in the multi-space decoder (14) may provide a solution to this problem by constraining $\hat{\mathbf{h}}_{\text{MS}}$ to belong to some prespecified feasible set. Fig. 1 gives an illustration of such a situation. The feasible set

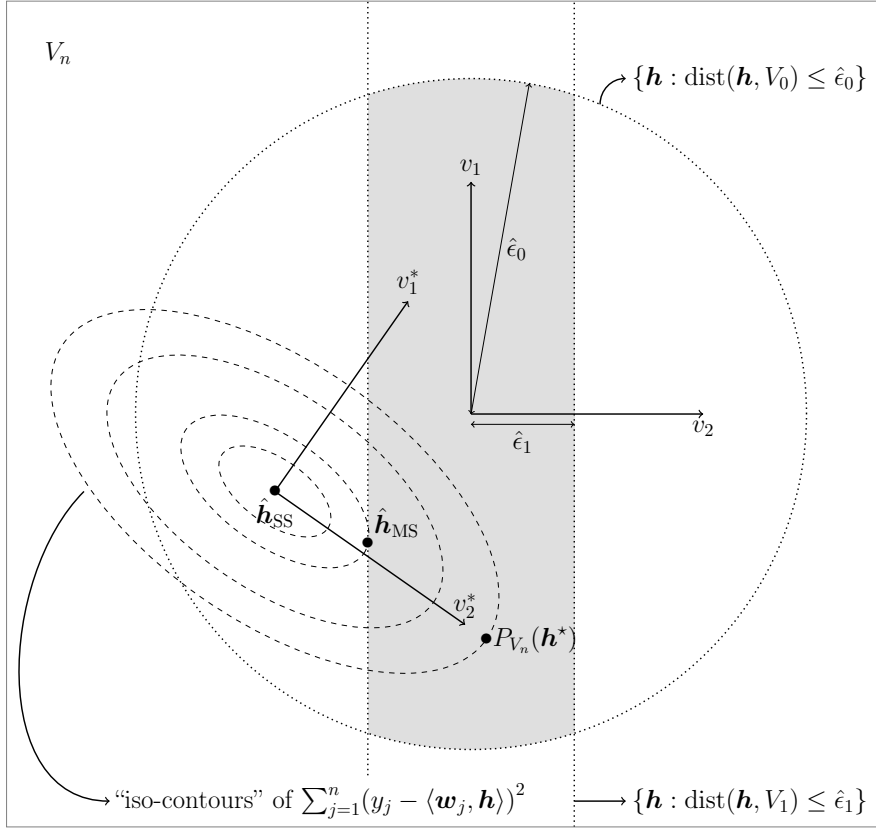


Fig. 1 Graphical representation of the geometry of the multi-space problem (14).

defined by the constraints in (14) is represented by the gray shaded area. In the simplified setup considered here, it corresponds to the intersection of two sets: the constraint associated to V_0 imposes $\hat{h}_{\text{MS}}$ to belong to a ball of radius $\hat{\epsilon}_0$; the constraint corresponding to V_1 requires that $\hat{h}_{\text{MS}}$ does not deviate from the line passing through v_1 by more than $\hat{\epsilon}_1$. The multi-space estimate $\hat{h}_{\text{MS}}$ then corresponds to the element of the feasible set leading to the smallest value of the cost function. We see in Fig. 1 that constraining the estimate $\hat{h}_{\text{MS}}$ to belong to the feasible set prevents it from deviating too far from $P_{V_n}(h^*)$. In particular, in the simple example described in Fig. 1, the multi-space estimate leads to better approximation performance than its single-space counterpart.

The same type of conclusions can in fact be drawn in more general setups: the multi-space decoder (14) is able to enhance the performance of the single-space approach (7) as soon as the prior information used in (14) can compensate for large deviations of the “iso-contour” ellipsoids. The gain achievable in the multi-space setup thus depends on the relative configuration of the “iso-contour ellipsoids” and the feasible set. We note that the shape of the ellipsoids

depends on the basis $\{\mathbf{v}_j^*\}_{j=1}^n$ and the singular value $\{\sigma_j\}_{j=1}^n$. Similarly, the feasible set is fully defined by the basis $\{\mathbf{v}_j\}_{j=1}^n$ and widths $\{\hat{\epsilon}_j\}_{j=0}^n$. Moreover, $\{\mathbf{v}_j\}_{j=1}^n$ and $\{\mathbf{v}_j^*\}_{j=1}^n$ only differ up to an orthogonal transformation $\mathbf{X}$, see (37). This explains why the parameters $\{\sigma_j\}_{j=1}^n$, $\{\hat{\epsilon}_j\}_{j=0}^n$ and $\{\delta_j\}_{j=1}^n$ (which depend on $\mathbf{X}$) play a crucial role in the characterization of the performance of the multi-space decoder in Theorem 2.

4 Proof of Theorem 2

In this section, we provide a proof of the result stated in Theorem 2. We first note that problem (14) is equivalent to finding the minimum of a quadratic function over a closed bounded subset of V_n . A minimizer thus always exists. Moreover, the unicity of the minimizer stated at the end of Theorem 2 follows from the strict convexity of the cost function over V_n when $\sigma_n > 0$.

In the rest of this section, we thus mainly focus on the derivation of the upper bound (21). Our proof is based on the following steps. First, since $\hat{\mathbf{h}}_{\text{MS}} \in V_n$, we have that

$$\begin{aligned} \|\mathbf{h}^* - \hat{\mathbf{h}}_{\text{MS}}\|^2 &= \|P_{V_n}(\mathbf{h}^*) - \hat{\mathbf{h}}_{\text{MS}}\|^2 + \|P_{V_n}^\perp(\mathbf{h}^*)\|^2, \\ &= \|P_{V_n}(\mathbf{h}^*) - \hat{\mathbf{h}}_{\text{MS}}\|^2 + \epsilon_n^2, \end{aligned} \quad (38)$$

where $P_{V_n}(\cdot)$ (resp. $P_{V_n}^\perp(\cdot)$) denotes the orthogonal projector onto V_n (resp. $V_n^\perp$). We then derive an upper bound on $\|P_{V_n}(\mathbf{h}^*) - \hat{\mathbf{h}}_{\text{MS}}\|^2$ as follows:

- We identify a set $\mathcal{D}$ such that $P_{V_n}(\mathbf{h}^*) - \hat{\mathbf{h}}_{\text{MS}} \in \mathcal{D}$ in Section 4.1. This implies in particular that $\|P_{V_n}(\mathbf{h}^*) - \hat{\mathbf{h}}_{\text{MS}}\|^2 \leq \sup_{\mathbf{d} \in \mathcal{D}} \|\mathbf{d}\|^2$.
- We derive the analytical expression of $\sup_{\mathbf{d} \in \mathcal{D}} \|\mathbf{d}\|^2$ as a function of the parameters $\{\epsilon_k\}_{k=0}^n$, $\{\hat{\epsilon}_k\}_{k=0}^n$ and $\{\sigma_k\}_{k=1}^n$ in Section 4.2.

Combining these results, we obtain (21)-(23).

4.1 Definition of $\mathcal{D}$

We express $\mathcal{D}$ as the intersection of two sets $\mathcal{D}_1$ and $\mathcal{D}_2$ that we define in Sections 4.1.2 and 4.1.3 respectively. In order to properly define these quantities, we introduce some particular orthonormal bases for V_n and $W_m = \text{span}(\{\mathbf{w}_j\}_{j=1}^m)$ in Section 4.1.1.

4.1.1 Some particular bases for V_n and W_m

Let

$$\mathbf{G} = \mathbf{U}\mathbf{A}\mathbf{X}^T \quad (39)$$

be the singular value decomposition of the Gram matrix defined in (12), where $\mathbf{U} \in \mathbb{R}^{m \times m}$ and $\mathbf{X} \in \mathbb{R}^{n \times n}$ are orthonormal matrices and $\mathbf{A} \in \mathbb{R}^{m \times n}$ is the diagonal matrix of singular values. We denote by $\{\sigma_j\}_{j=1}^n$ the set of singular values of $\mathbf{G}$ sorted in their decreasing order of magnitude. We remind the reader that if $m < n$, we adopt the convention $\sigma_j = 0 \ \forall j > m$.

We define the following bases for V_n and W_m :

$$\mathbf{v}_j^* = \sum_{i=1}^n x_{ij} \mathbf{v}_i, \quad (40)$$

$$\mathbf{w}_j^* = \sum_{i=1}^m u_{ij} \mathbf{w}_i, \quad (41)$$

where $\mathbf{U} \in \mathbb{R}^{m \times m}$ and $\mathbf{X} \in \mathbb{R}^{n \times n}$ are the orthonormal matrices appearing in (39). We note that $\{\mathbf{v}_j^*\}_{j=1}^n$ is an orthonormal basis whereas $\{\mathbf{w}_j^*\}_{j=1}^m$ is not necessarily orthonormal. By definition, $\{\mathbf{v}_j^*\}_{j=1}^n$ and $\{\mathbf{w}_j^*\}_{j=1}^m$ enjoy the following desirable property:

$$\langle \mathbf{w}_i^*, \mathbf{v}_j^* \rangle = \begin{cases} \sigma_j & \text{if } i = j \\ 0 & \text{otherwise.} \end{cases} \quad (42)$$

4.1.2 Definition of $\mathcal{D}_1$

Let us define $\mathcal{D}_1$ as

$$\mathcal{D}_1 = \left\{ \mathbf{d} = \sum_{j=1}^n \beta_j \mathbf{v}_j^* : \sum_{j=1}^n \sigma_j^2 \beta_j^2 \leq 4\gamma^2 \epsilon_n^2 \right\}, \quad (43)$$

where γ is defined in (18). We show hereafter that $P_{V_n}(\mathbf{h}^*) - \hat{\mathbf{h}}_{\text{MS}} \in \mathcal{D}_1$.

Let us first consider the intermediate set

$$\mathcal{S} = \{\mathbf{h} : f(\mathbf{h}) \leq \gamma^2 \epsilon_n^2\}, \quad (44)$$

where $f(\mathbf{h}) \triangleq \sum_{j=1}^m (y_j - \langle \mathbf{w}_j, \mathbf{h} \rangle)^2$ is the cost function appearing in the variational formulation of multi-space decoder (14).

Clearly $P_{V_n}(\mathbf{h}^*) \in \mathcal{S}$ because

$$\begin{aligned}
 f(P_{V_n}(\mathbf{h}^*)) &= \sum_{j=1}^m (y_j - \langle \mathbf{w}_j, P_{V_n}(\mathbf{h}^*) \rangle)^2 \\
 &= \sum_{j=1}^m (\langle \mathbf{w}_j, \mathbf{h}^* \rangle - \langle \mathbf{w}_j, P_{V_n}(\mathbf{h}^*) \rangle)^2 \\
 &= \sum_{j=1}^m (\langle \mathbf{w}_j, P_{V_n}^\perp(\mathbf{h}^*) \rangle)^2 \\
 &\leq \gamma^2 \|P_{V_n}^\perp(\mathbf{h}^*)\|^2 \\
 &\leq \gamma^2 \epsilon_n^2.
 \end{aligned} \tag{45}$$

Moreover, we also have $\hat{\mathbf{h}}_{\text{MS}} \in \mathcal{S}$. This can be seen from the following arguments. First, $P_{V_n}(\mathbf{h}^*)$ is a feasible point for problem (14), that is

$$\text{dist}(P_{V_n}(\mathbf{h}^*), V_k) \leq \hat{\epsilon}_k \text{ for } k = 0 \dots n. \tag{46}$$

Indeed, rewriting $\mathbf{h}^*$ as

$$\mathbf{h}^* = \sum_{j=1}^n \langle \mathbf{v}_j, \mathbf{h}^* \rangle \mathbf{v}_j + \mathbf{z}, \tag{47}$$

where $\mathbf{z} \in V_n^\perp$, we have

$$\begin{aligned}
 \hat{\epsilon}_k &\geq \text{dist}(\mathbf{h}^*, V_k) \\
 &= \|P_{V_k}^\perp(\mathbf{h}^*)\| \\
 &= \left\| \sum_{j=k+1}^n \langle \mathbf{v}_j, \mathbf{h}^* \rangle \mathbf{v}_j + \mathbf{z} \right\| \\
 &= \sqrt{\left\| \sum_{j=k+1}^n \langle \mathbf{v}_j, \mathbf{h}^* \rangle \mathbf{v}_j \right\|^2 + \|\mathbf{z}\|^2} \\
 &\geq \left\| \sum_{j=k+1}^n \langle \mathbf{v}_j, \mathbf{h}^* \rangle \mathbf{v}_j \right\| \\
 &= \|P_{V_k}^\perp(P_{V_n}(\mathbf{h}^*))\| \\
 &= \text{dist}(P_{V_n}(\mathbf{h}^*), V_k).
 \end{aligned} \tag{48}$$

The first inequality follows from our initial assumption (16). The third equality is true because $\mathbf{z} \in V_n^\perp$. Now, since $\hat{\mathbf{h}}_{\text{MS}}$ is a minimizer of $f(\mathbf{h})$ over the set of feasible points, we have $f(\hat{\mathbf{h}}_{\text{MS}}) \leq f(P_{V_n}(\mathbf{h}^*)) \leq \gamma^2 \epsilon_n^2$ and therefore $\hat{\mathbf{h}}_{\text{MS}} \in \mathcal{S}$.

Finally, we show that $\hat{\mathbf{h}}_{\text{MS}} \in \mathcal{S}$ and $P_{V_n}(\mathbf{h}^*) \in \mathcal{S}$ implies $P_{V_n}(\mathbf{h}^*) - \hat{\mathbf{h}}_{\text{MS}} \in \mathcal{D}_1$. Let us first note that, if $\mathbf{h} \in V_n$, the cost function $f(\mathbf{h})$ can be rewritten as:

$$\begin{aligned} f(\mathbf{h}) &= \sum_{j=1}^m (\langle \mathbf{w}_j, \mathbf{h}^* \rangle - \langle \mathbf{w}_j, \mathbf{h} \rangle)^2, \\ &= \sum_{j=1}^m (\langle \mathbf{w}_j^*, \mathbf{h}^* \rangle - \langle \mathbf{w}_j^*, \mathbf{h} \rangle)^2, \\ &= \sum_{j=1}^n (\langle \mathbf{w}_j^*, \mathbf{h}^* \rangle - \sigma_j \langle \mathbf{v}_j^*, \mathbf{h} \rangle)^2 + \sum_{j=n+1}^m \langle \mathbf{w}_j^*, \mathbf{h}^* \rangle^2, \end{aligned} \quad (49)$$

where the second equality follows from the fact that $\{\mathbf{w}_j\}_{j=1}^m$ and $\{\mathbf{w}_j^*\}_{j=1}^m$ differ up to an orthonormal transformation; the last equality is a consequence of (42) and the fact that $\mathbf{h} \in V_n$ by hypothesis.

We note that, since $\{\mathbf{v}_j^*\}_{j=1}^n$ is an orthonormal basis of V_n , $P_{V_n}(\mathbf{h}^*) - \hat{\mathbf{h}}_{\text{MS}}$ can be written as $\sum_{j=1}^n \beta_j \mathbf{v}_j^*$ by setting $\beta_j = \langle \mathbf{v}_j^*, P_{V_n}(\mathbf{h}^*) \rangle - \langle \mathbf{v}_j^*, \hat{\mathbf{h}}_{\text{MS}} \rangle$. Therefore, we have

$$\begin{aligned} \sum_{j=1}^n \sigma_j^2 \beta_j^2 &= \sum_{j=1}^n \left(\sigma_j \langle \mathbf{v}_j^*, P_{V_n}(\mathbf{h}^*) \rangle - \sigma_j \langle \mathbf{v}_j^*, \hat{\mathbf{h}}_{\text{MS}} \rangle \right)^2, \\ &= \sum_{j=1}^n \left(\sigma_j \langle \mathbf{v}_j^*, P_{V_n}(\mathbf{h}^*) \rangle - \langle \mathbf{w}_j^*, \mathbf{h}^* \rangle - \sigma_j \langle \mathbf{v}_j^*, \hat{\mathbf{h}}_{\text{MS}} \rangle + \langle \mathbf{w}_j^*, \mathbf{h}^* \rangle \right)^2, \\ &\leq 2 \sum_{j=1}^n \left(\sigma_j \langle \mathbf{v}_j^*, P_{V_n}(\mathbf{h}^*) \rangle - \langle \mathbf{w}_j^*, \mathbf{h}^* \rangle \right)^2 + 2 \sum_{j=1}^n \left(\sigma_j \langle \mathbf{v}_j^*, \hat{\mathbf{h}}_{\text{MS}} \rangle - \langle \mathbf{w}_j^*, \mathbf{h}^* \rangle \right)^2, \\ &\leq 2f(P_{V_n}(\mathbf{h}^*)) + 2f(\hat{\mathbf{h}}_{\text{MS}}), \\ &\leq 4\gamma^2 \epsilon_n^2, \end{aligned}$$

where the first inequality follows from the standard inequality $(a + b)^2 \leq 2(a^2 + b^2)$, the second from (49), and the last one from the fact that $\hat{\mathbf{h}}_{\text{MS}} \in \mathcal{S}$ and $P_{V_n}(\mathbf{h}^*) \in \mathcal{S}$.

4.1.3 Definition of $\mathcal{D}_2$

Let

$$\delta_j = \eta_j + \hat{\eta}_j, \quad (50)$$

where

$$\begin{aligned}\eta_j &= \sum_{i=1}^n |x_{ij}| \epsilon_{i-1}, \\ \hat{\eta}_j &= \sum_{i=1}^n |x_{ij}| \hat{\epsilon}_{i-1},\end{aligned}\tag{51}$$

and the x_{ij} 's are the elements of the matrix $\mathbf{X}$ appearing in the SVD decomposition (39). We define $\mathcal{D}_2$ as

$$\mathcal{D}_2 = \left\{ \mathbf{d} = \sum_{j=1}^n \beta_j \mathbf{v}_j^* : |\beta_j| \leq \eta_j \right\}.\tag{52}$$

We show hereafter that $P_{V_n}(\mathbf{h}^*) - \hat{\mathbf{h}}_{\text{MS}} \in \mathcal{D}_2$.

We first note that if $\mathbf{h}$ is feasible for problem (14), we must have

$$|\langle \mathbf{v}_j^*, \mathbf{h} \rangle| \leq \hat{\eta}_j.\tag{53}$$

Indeed, if $\mathbf{h}$ is feasible, the constraint $\text{dist}(\mathbf{h}, V_k) \leq \hat{\epsilon}_k$ simply writes as

$$\sum_{j=k+1}^n \langle \mathbf{v}_j, \mathbf{h} \rangle^2 \leq \hat{\epsilon}_k^2.$$

In particular, this implies that

$$|\langle \mathbf{v}_{k+1}, \mathbf{h} \rangle| \leq \hat{\epsilon}_k.$$

Using the fact that

$$\mathbf{v}_j^* = \sum_{k=1}^n x_{kj} \mathbf{v}_k,$$

we obtain (53). In a similar way, we can find that

$$|\langle \mathbf{v}_j^*, P_{V_n}(\mathbf{h}^*) \rangle| \leq \eta_j,\tag{54}$$

by using the fact that $\text{dist}(P_{V_n}(\mathbf{h}^*), V_k) \leq \epsilon_k$ from (48).

Let us now show that $P_{V_n}(\mathbf{h}^*) - \hat{\mathbf{h}}_{\text{MS}} \in \mathcal{D}_2$. Since $\{\mathbf{v}_j^*\}_{j=1}^n$ is an orthonormal basis of V_n , $P_{V_n}(\mathbf{h}^*) - \hat{\mathbf{h}}_{\text{MS}}$ can be written as $\sum_{j=1}^n \beta_j \mathbf{v}_j^*$ by setting $\beta_j = \langle \mathbf{v}_j^*, P_{V_n}(\mathbf{h}^*) \rangle - \langle \mathbf{v}_j^*, \hat{\mathbf{h}}_{\text{MS}} \rangle$. This leads to

$$\begin{aligned}|\beta_j| &= \left| \langle \mathbf{v}_j^*, P_{V_n}(\mathbf{h}^*) \rangle - \langle \mathbf{v}_j^*, \hat{\mathbf{h}}_{\text{MS}} \rangle \right|, \\ &\leq \left| \langle \mathbf{v}_j^*, P_{V_n}(\mathbf{h}^*) \rangle \right| + \left| \langle \mathbf{v}_j^*, \hat{\mathbf{h}}_{\text{MS}} \rangle \right|, \\ &\leq \hat{\eta}_j + \eta_j = \delta_j,\end{aligned}$$

where the last inequality follows from (53) and (54).

4.2 Expression of $\sup_{\mathbf{d} \in \mathcal{D}} \|\mathbf{d}\|^2$

We consider the following problem:

$$\sup_{\mathbf{d} \in \mathcal{D}} \|\mathbf{d}\|^2 = \sup_{\boldsymbol{\beta}} \|\boldsymbol{\beta}\|^2 \quad \text{subject to} \quad \begin{cases} \sum_{j=1}^n \sigma_j^2 \beta_j^2 \leq 4\gamma^2 \epsilon_n^2 \\ |\beta_j| \leq \delta_j \end{cases}. \quad (55)$$

If $\sum_{j=1}^n \sigma_j^2 \delta_j^2 < 4\gamma^2 \epsilon_n^2$, the first constraint in (55) is always inactive and the solution simply reads

$$\sup_{\mathbf{d} \in \mathcal{D}} \|\mathbf{d}\|^2 = \sum_{j=1}^n \delta_j^2. \quad (56)$$

If $\sum_{j=1}^n \sigma_j^2 \delta_j^2 \geq 4\gamma^2 \epsilon_n^2$, the solution of (55) is given by

$$\sup_{\mathbf{d} \in \mathcal{D}} \|\mathbf{d}\|^2 = \sum_{j=\ell+1}^n \delta_j^2 + \rho \delta_\ell^2, \quad (57)$$

where ℓ is the largest integer such that

$$\sum_{j=\ell}^n \sigma_j^2 \delta_j^2 \geq 4\gamma^2 \epsilon_n^2, \quad (58)$$

and $\rho \in [0, 1]$ is defined as

$$\rho \sigma_\ell^2 \delta_\ell^2 + \sum_{j=\ell+1}^n \sigma_j^2 \delta_j^2 = 4\gamma^2 \epsilon_n^2. \quad (59)$$

This can be seen by verifying the optimality condition of problem (55). We note that problem (55) is the same (up to some constants) to the one considered in [5, Section 3.1]. The solution (57) is therefore similar, up to some different constants, to the one obtained in that paper.

5 Conclusions

In this paper, we provide a mathematical characterization of the performance of some particular decoder exploiting a set of approximation subspaces. This decoder was previously shown to lead to good empirical results in several contributions [1, 8–10], although no proof of its theoretical performance was provided in these works. Our result shows how the performance of this “multi-space” decoder is related to the parameters defining the approximation problem: the observation W_m and prior subspaces $\{V_k\}_{k=0}^n$, the quality of the prior information (*i.e.*, the widths $\{\hat{\epsilon}_k\}_{k=0}^n$) and the distance between the target solution $\mathbf{h}^*$ and the approximation subspaces (*i.e.*, $\{\epsilon_k\}_{k=0}^n$). Based on this result, we show that the “multi-space” decoder can have provably better reconstruction guarantees than its standard (“single-space”) counterpart in some situations.

A Proof of (36)

In this appendix, we show that the cost function $f(\mathbf{h}) \triangleq \sum_{j=1}^n (y_j - \langle \mathbf{w}_j, \mathbf{h} \rangle)^2$ can be rewritten as in (36) when $\mathbf{h} \in V_n$ and $\sigma_n > 0$.⁶ First, using the definition of y_j , we have

$$f(\mathbf{h}) = \sum_{j=1}^n (\langle \mathbf{w}_j, \mathbf{h}^* \rangle - \langle \mathbf{w}_j, \mathbf{h} \rangle)^2. \quad (60)$$

Moreover, using the particular orthonormal bases introduced in Section 4.1.1, we obtain

$$\begin{aligned} f(\mathbf{h}) &= \sum_{j=1}^n (\langle \mathbf{w}_j^*, \mathbf{h}^* \rangle - \langle \mathbf{w}_j^*, \mathbf{h} \rangle)^2, \\ &= \sum_{j=1}^n (\langle \mathbf{w}_j^*, \mathbf{h}^* \rangle - \sigma_j \langle \mathbf{v}_j^*, \mathbf{h} \rangle)^2, \end{aligned} \quad (61)$$

where the first equality follows from the fact that $\{\mathbf{w}_j\}_{j=1}^n$ and $\{\mathbf{w}_j^*\}_{j=1}^n$ differ up to an orthogonal transformation; the second is a consequence of (42) and our hypothesis $\mathbf{h} \in V_n$.

Since $\hat{\mathbf{h}}_{\text{SS}}$ corresponds to the minimum of $f(\mathbf{h})$ over V_n (see (8)), we simply have

$$\langle \mathbf{v}_j^*, \hat{\mathbf{h}}_{\text{SS}} \rangle = \frac{\langle \mathbf{w}_j^*, \mathbf{h}^* \rangle}{\sigma_j}, \quad (62)$$

if $\sigma_n > 0$. Hence, under this assumption, (61) can also be rewritten as in (36).

References

1. Argaud, J., Bouriquet, B., Gong, H., Maday, Y., Mula, O.: Stabilization of (G)EIM in presence of measurement noise: application to nuclear reactor physics. In: Spectral and High Order Methods for Partial Differential Equations ICOSAHOM 2016, pp. 133–145. Springer (2017)
2. Babuska, I.: Error-bounds for finite element method. *Numerische Mathematik* **16**, 322–333 (1970/71)
3. Barrault, M., Maday, Y., Nguyen, N.C., Patera, A.T.: An ‘empirical interpolation’ method: application to efficient reduced-basis discretization of partial differential equations. *Comptes Rendus Mathématique* **339**(9), 667–672 (2004). DOI 10.1016/j.crma.2004.08.006. URL <http://www.sciencedirect.com/science/article/pii/S1631073X04004248>
4. Berkooz, G., Holmes, P., Lumley, J.L.: The proper orthogonal decomposition in the analysis of turbulent flows. *Annual Review of Fluid Mechanics* **25**(1), 539–575 (1993). DOI 10.1146/annurev.fl.25.010193.002543. URL <http://dx.doi.org/10.1146/annurev.fl.25.010193.002543>
5. Binev, P., Cohen, A., Dahmen, W., DeVore, R., Petrova, G., Wojtaszczyk, P.: Data assimilation in reduced modeling. *SIAM/ASA Journal on Uncertainty Quantification* **5**(1), 1–29 (2017). DOI 10.1137/15M1025384. URL <https://doi.org/10.1137/15M1025384>
6. Chaturantabut, S., Sorensen, D.: Nonlinear model reduction via discrete empirical interpolation. *SIAM J. Sci. Comput.* **32**(5), 2737–2764 (2010). DOI 10.1137/090766498. URL <http://dx.doi.org/10.1137/090766498>
7. Everson, R., Sirovich, L.: Karhunen-Loève procedure for gappy data. *J. Opt. Soc. Am. A* **12**(8), 1657–1664 (1995). DOI 10.1364/josaa.12.001657. URL <http://dx.doi.org/10.1364/josaa.12.001657>

⁶ We remind the reader that we assume $m = n$.

8. Fick, L., Maday, Y., Patera, A.T., Taddei, T.: A Reduced Basis Technique for Long-Time Unsteady Turbulent Flows. ArXiv e-prints (2017)
9. Herzet, C., Diallo, M., Héas, P.: Beyond petrov-galerkin projection by using multi-space prior. In: European Conference on Numerical Mathematics and Advanced Applications (Enumath'17). Voss, Norway (2017)
10. Herzet, C., Diallo, M., Héas, P.: Beyond Petrov-Galerkin projection by using multi-space prior. In: Model Reduction of Parametrized Systems IV (MoRePaS'18). Nantes, France (2018)
11. Quarteroni, A., Manzoni, A., Negri, F.: Reduced Basis Methods for Partial Differential Equations, vol. 92. Springer International Publishing (2016). URL <http://www.springer.com/us/book/9783319154305>
12. Xu, J., Xikatanov, L.: Some observation on Babuska and Brezzi theories. Numer. Math. **94**(1), 195–202 (2003)