

Universität Stuttgart



B. Haasdonk

Reduced Basis Methods for Parametrized PDEs – A Tutorial Introduction for Stationary and Instationary Problems

Stuttgart, March 2014, revised version March 2016 *

Institute of Applied Analysis and Numerical Simulation,
University of Stuttgart
Pfaffenwaldring 57, 70569 Stuttgart / Germany
haasdonk@mathematik.uni-stuttgart.de
<http://www.ians.uni-stuttgart.de/agh>

Abstract In this part we are concerned with a class of model reduction techniques for parametric partial differential equations, the so-called Reduced Basis (RB) methods. These allow to obtain low-dimensional parametric models for various complex applications, enabling accurate and rapid numerical simulations. Important aspects are basis generation and certification of the simulation results by suitable a posteriori error control. The main terminology, ideas and assumptions will be explained for the case of linear stationary elliptic, as well as parabolic or hyperbolic instationary problems. Reproducible experiments will illustrate the theoretical findings. We close with a discussion of further recent developments.

Keywords Reduced Basis Methods, Parametrized Partial Differential Equations

* To appear as chapter in: P. Benner, A. Cohen, M. Ohlberger, K. Willcox: “Model Reduction and Approximation: Theory and Algorithms”, SIAM, Philadelphia, 2016.

Preprint Series
Stuttgart Research Centre for Simulation Technology (SRC SimTech)

Issue No. 2014

SimTech – Cluster of Excellence
Pfaffenwaldring 7a
70569 Stuttgart
publications@simtech.uni-stuttgart.de
www.simtech.uni-stuttgart.de

Reduced Basis Methods for Parametrized PDEs – A Tutorial Introduction for Stationary and Instationary Problems

Bernard Haasdonk
University of Stuttgart
Institute of Applied Analysis and Numerical Simulation
Pfaffenwaldring 57
70569 Stuttgart
Germany
haasdonk@mathematik.uni-stuttgart.de

Abstract

In this part we are concerned with a class of model reduction techniques for parametric partial differential equations, the so-called Reduced Basis (RB) methods. These allow to obtain low-dimensional parametric models for various complex applications, enabling accurate and rapid numerical simulations. Important aspects are basis generation and certification of the simulation results by suitable a posteriori error control. The main terminology, ideas and assumptions will be explained for the case of linear stationary elliptic, as well as parabolic or hyperbolic instationary problems. Reproducible experiments will illustrate the theoretical findings. We close with a discussion of further recent developments.

1 Introduction

Discretization techniques for partial differential equations (PDEs) frequently lead to very high-dimensional numerical models with corresponding high demands concerning hardware and computation times. This is the case for various discretization types, such as Finite Elements (FE), Finite Volumes (FV), Discontinuous Galerkin (DG) methods, etc. These high computational costs pose a serious problem in the context of multi-query, real-time or slim computing scenarios: *Multi-query* scenarios comprise settings, where not one single simulation result is required, but the setting of the problem is varying, and multiple simulation requests are requested. Such situations can be observed in the case in parameter studies, design, optimization, inverse problems or statistical analysis. *Real-time* scenarios consist of problems, where the simulation result is required very fast. This can be the case for simulation-based interaction with real processes, e.g. control or prediction, or interaction with humans,

e.g. a development engineer working with simulation software and requiring rapid answers. *Slim computing* scenarios denote settings, where computational capabilities are very limited with respect to speed or memory but still accurate simulation answers are required. This can comprise simple technical controllers, smartphone apps, etc.

In the above scenarios, the “varying” quantities which describe the problem, will be denoted as *parameters* and are collected in a parameter vector $\mu \in \mathcal{P}$. Here we assume that $\mathcal{P} \subset \mathbb{R}^p$ is a set of possible/admitted parameters of low dimension p . The parametric solution then will be denoted $u(\mu)$ and typically stems from a solution space X that can be infinite or at least very high-dimensional. Frequently, not the solution itself, but rather a quantity of interest $s(\mu)$ depending on the solution is desired. So, the standard computational procedure is to start with a low-dimensional parameter, compute a typically high-dimensional solution $u(\mu)$ and then derive a low-dimensional output quantity. Clearly, the computationally intensive part in this chain is the computation of the solution $u(\mu)$. Therefore, the aim of *model reduction* is to develop techniques that provide low-dimensional and hence rapidly computable approximations for the solution $u(\mu)$ and possible outputs. *Reduced Basis* (RB) methods focus on a certain problem class, namely parametric PDEs. The crucial insight enabling simplification of parametric problems is the fact that the solution manifold $\mathcal{M}$, i.e. the set of parametric solutions, often can be well approximated by a low-dimensional subspace $X_N \subset X$. In RB-methods one popular way is the construction of this subspace by *snapshots*, i.e. X_N is spanned by solutions $u(\mu^{(i)})$ for suitable parameters $\mu^{(i)}, i = 1, \dots, N$. A crucial requirement for constructing a good approximating space is a careful choice of these parameters. After construction of the space, the reduced model is obtained, e.g. by Galerkin projection, and provides an approximation $u_N(\mu) \in X_N$ of the solution and an approximation $s_N(\mu)$ of the output quantity of interest. See Fig. 1 for an illustration of the RB-approximation scenario. In addition to the pure reduction, also error control is desired, i.e. avail-

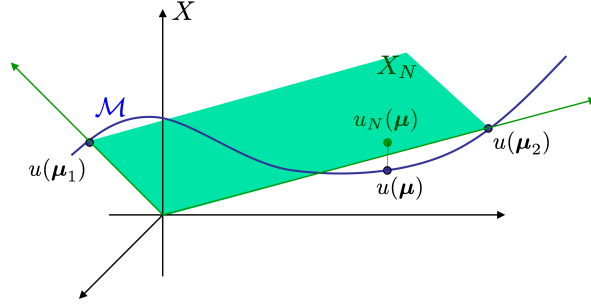


Figure 1: Illustration of the solution manifold and the RB-approximation.

ability of computable and *rigorous*, that means provable, upper bounds for the state or output error. Additionally, these error bounds should be *effective*, i.e. not arbitrarily overestimate the error. The computational procedure is ideally decomposed

in an offline and online phase: During the *offline* phase, performed once, a reduced basis is generated and further auxiliary quantities are precomputed. Then, in the *online* phase, for varying parameters μ , the approximate solution, output and error bounds can be provided rapidly. The computational complexity of the online phase will usually not depend on the dimension of the full space X , hence, the space X can be assumed to be arbitrarily accurate. Instead, the computational complexity of the online phase will typically be only polynomial in N , the dimension of the reduced space. The runtime advantage of an RB-model in the context of a multi-query scenario is illustrated in Fig. 2: The offline phase is typically much more expensive than

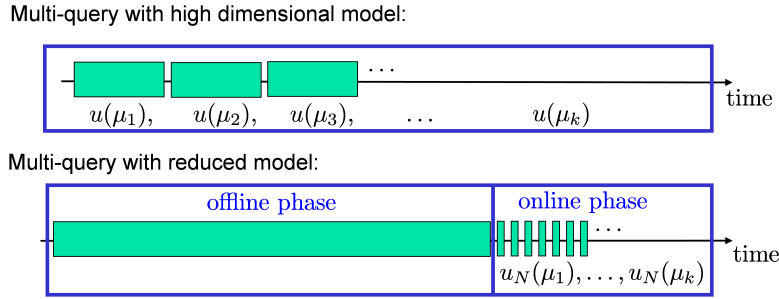


Figure 2: Runtime advantage of RB-model in multi-query scenarios.

several full simulations. However, if a sufficient number of reduced solutions are required in the online phase, the overall computation time will be decreased in contrast to many full simulations. We collect some of the motivating questions that will be addressed in this tutorial:

- How can we construct good spaces X_N ? Can such procedures be provably “good”?
- How can we obtain a good approximation $u_N(\mu) \in X_N$?
- How can $u_N(\mu)$ be determined rapidly, i.e. computationally efficient?
- Can stability or convergence with growing N be obtained?
- Can the RB-error be rigorously bounded? Are the error bounds fully computable?
- Are the error bounds largely overestimating the error, or can the “effectivity” be quantified?
- For which problem classes can low-dimensional approximation expected to be successful?

RB-methods can be traced back to the last century [21, 55, 58, 60] but received plenty of attention in the last decade in terms of *certification*, i.e. providing error and

accuracy control, leading to applicable and efficient procedures for various problem types. Problems such as stationary linear elliptic, stationary non-coercive, saddle point problems in particular Stokes flow, instationary parabolic and hyperbolic and nonlinear problems have been treated, geometry parametrizations can be handled, etc. Various papers and PhD theses have meanwhile been devoted to the topic. For the moment, we refer to the electronic book [57], the overview article [63] and references therein. Also, we want to refer to the excellent collection of papers and theses at <http://augustine.mit.edu>. Concerning software, also different packages have been developed, which address RB-methods. Apart from our package *RBmatlab*, available for download at the website www.morepas.org/software, we mention the packages *rbMIT* and *rbAPPmit*, available at augustine.mit.edu/methodology and *pymor*, publically accessible at the website github.com/pymor. Further references to papers and other electronic resources can be found at www.morepas.org and www.modelreduction.org. We postpone giving individual further references to the concluding Sec. 4. The purpose of the present chapter is to provide a tutorial introduction to RB-methods. It may serve as material for an introductory course on the topic. The suggested audience are students or researchers that have a background in numerical analysis for PDEs and elementary discretization techniques. We aim at a self-contained document, collecting the central statements and providing elementary proofs or explicit references to corresponding literature. We include experiments that can be reproduced with the software package *RBmatlab*. We also provide plenty of exercises that are recommended for deepening the theoretical understanding of the present methodology.

The document is structured into two main parts. The first part consists of Sec. 2, where we consider elliptic coercive problems. This material is largely based on existing work, in particular [57] and lecture notes [27]. We devote the second part, Sec. 3, to the time-dependent case and give corresponding RB-formulations. We close this chapter by providing an outlook and references on further topics and recent developments in Sec. 4. A selection of accompanying exercises is given in Appendix 5.

2 Stationary Problems

We start with stationary problems and focus on symmetric or nonsymmetric elliptic partial differential equations. Overall, the collection of results in this chapter serves as a *general RB-pattern* for new problem classes. This means that the current sequence of results/procedures can be used as a schedule for other problems. One can try to sequentially obtain analogous results for a new problem along the lines of this section.

Note, that the results and procedures of this section are mostly well known and can be considered to be standard. Hence, we do not claim any (major) novelty in the current section, but rather see it as a collection and reformulation of existing results and methodology. We introduce slight extensions or intermediate results at some points. Some references that must be attributed are [57, 63] and references therein, but similar formulations also can be found in further publications.

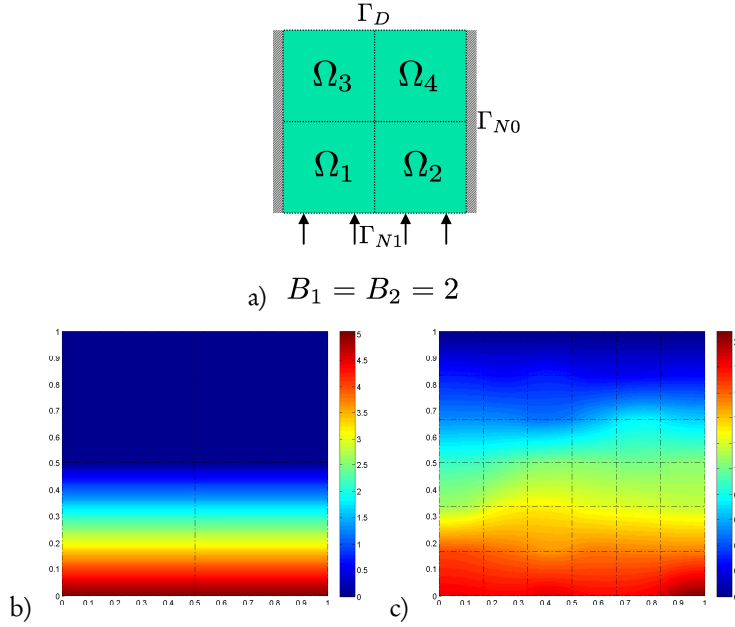


Figure 3: Thermal block: a) geometry and notation, b) solution sample for $B_1 = B_2 = 2$ and $\mu = (1, 1, 1, 1)^T$, c) solution sample for $B_1 = B_2 = 6$ and random parameter μ .

2.1 Model Problem

A very elegant model problem has been given in [57], which we also want to adopt (with minor modification) as a driving model example for the methodology in this section. It is an example of a parametrized partial differential equation modelling the heat transport through a block of solid material that is assembled by subblocks of different heat conductivities. The values of the piecewise constant heat conductivities are considered as parameters in the problem. Consequently, the example is called a thermal block. The block is heated on a part of its boundary, insulated on other parts and cooled to a reference temperature on the remaining boundary part. We are interested in the average temperature on the heating boundary part.

Fig. 3a) explains the geometry and the notation: Let $\Omega = (0, 1)^2$ be the unit square and $B_1, B_2 \in \mathbb{N}$ the number of subblocks per dimension. The subblocks are denoted $\Omega_i, i = 1, \dots, p$ for $p := B_1 B_2$ counted rowwise starting from the left bottom. The bottom boundary is denoted $\Gamma_{N,1}$ with unit outward normal $n(x)$, where we will prescribe a unit flux into the domain. The left and right boundary are insulated no-flux boundaries denoted by $\Gamma_{N,0}$. The upper boundary Γ_D is a homogeneous Dirichlet boundary, where we assign 0 as temperature. The heat conductivities are defined as parameters $\mu_i, i = 1, \dots, p$. We prescribe a suitable parameter domain for the parameter vector $\mu = (\mu_1, \dots, \mu_p)^T \in \mathcal{P} := [\mu_{\min}, \mu_{\max}]^p$, namely logarithmically symmetric around 1, i.e. $\mu_{\min} = 1/\mu_{\max}$ for $\mu_{\max} > 1$. Note, that the model

in [57] further assumes the first parameter to be normalized to 1, hence that formulation essentially has one parameter less than the current formulation. The space- and parameter-dependent heat-conductivity function then is expressed as piecewise constant function via indicator functions χ_{Ω_q}

$$\kappa(x; \mu) := \sum_{q=1}^p \mu_q \chi_{\Omega_q}(x).$$

The parametrized PDE that needs to be solved for the parametric solution $u(x; \mu)$ is the elliptic problem

$$\begin{aligned} -\nabla \cdot (\kappa(x; \mu) \nabla u(x; \mu)) &= 0, & x \in \Omega \\ u(x; \mu) &= 0, & x \in \Gamma_D \\ (\kappa(x; \mu) \nabla u(x; \mu)) \cdot n(x) &= i, & x \in \Gamma_{N,i}, i = 0, 1. \end{aligned}$$

The weak form is based on the solution space $H_{\Gamma_D}^1(\Omega) := \{u \in H^1(\Omega) \mid u|_{\Gamma_D} = 0\}$ of functions vanishing on Γ_D (in the trace sense). Here $H^1(\Omega)$ denotes the standard Sobolev space of square integrable functions, which have square integrable derivatives. Then, for given $\mu \in \mathcal{P}$ we are interested in the solution $u(\cdot; \mu) \in H_{\Gamma_D}^1(\Omega)$ such that

$$\sum_{q=1}^p \int_{\Omega_q} \mu_q \nabla u(x; \mu) \cdot \nabla v(x) dx = \int_{\Gamma_{N,1}} v(x) dx$$

for all test functions $v \in H_{\Gamma_D}^1(\Omega)$. Then, we evaluate a scalar output value, e.g. the average temperature at the bottom

$$s(\mu) := \int_{\Gamma_{N,1}} u(x; \mu) dx.$$

This model allows some simple but interesting insights into the structure of the solution manifold for varying μ .

- Simple solution structure: In the case of $B_1 = 1$ (or $B_1 > 1$ but identical parameters in each row), the solution exhibits horizontal symmetry, cf. Fig. 3b). One can easily show that the solution for all $\mu \in \mathcal{P}$ is piecewise linear and contained in a B_2 -dimensional linear subspace of the infinite-dimensional space $H_{\Gamma_D}^1(\Omega)$, cf. Exercise 5.1.
- Complex solution structure: Plot c) indicates a more complex example, where the solution manifold for $B_1 > 1$ with independent parameters cannot be exactly approximated by a finite-dimensional solution space.
- Parameter redundancy: The solution manifold is invariant with respect to scaling of the parameter vector, i.e. if $u(\mu)$ is a given solution, then $u(c\mu) = \frac{1}{c}u(\mu)$ for $c > 0$ is the solution for the parameter $c\mu$. This is an important insight for parametric models: More parameters do not necessarily increase the parametric complexity of the solution manifold.

The thermal block is an excellent motivation and opportunity for RB-methods: Solution manifolds may possess structural simplicity or redundancy although the solution space is very high- or even infinite-dimensional. Identifying such structures in the solution manifold may offer chances for low-dimensional accurate approximation. Based on this example of a parametrized partial differential equation, we can formulate the abstract setting.

2.2 Full Problem

An abstract formulation for a large class of linear stationary problems will be given. This will be the basis for the exposition in the subsequent sections. We assume X to be a real, separable Hilbert space with inner product $\langle \cdot, \cdot \rangle$, norm $\|\cdot\|$ and dual space X' with norm $\|\cdot\|_{X'}$. We assume to have a parameter domain $\mathcal{P} \subset \mathbb{R}^p$, a parameter-dependent bilinear form $a(\cdot, \cdot; \mu)$ and linear forms $l(\cdot; \mu), f(\cdot; \mu) \in X'$ for all $\mu \in \mathcal{P}$. We do not require symmetry for $a(\cdot, \cdot; \mu)$. We assume the bilinear form and the linear forms to be uniformly continuous and the bilinear form to be uniformly coercive in the following sense:

Definition 2.1 (Uniform Continuity and Coercivity). *The parametric bilinear form $a(\cdot, \cdot; \mu)$ is assumed to be continuous, i.e. there exists $\gamma(\mu) \in \mathbb{R}$ with*

$$\gamma(\mu) := \sup_{u, v \in X \setminus \{0\}} \frac{a(u, v; \mu)}{\|u\| \|v\|} < \infty$$

and the continuity is uniform with respect to μ in the sense that for some $\bar{\gamma} < \infty$ holds $\gamma(\mu) \leq \bar{\gamma}$ for all $\mu \in \mathcal{P}$. Further, $a(\cdot, \cdot; \mu)$ is assumed to be coercive, i.e. there exists $\alpha(\mu)$ with

$$\alpha(\mu) := \inf_{u \in X \setminus \{0\}} \frac{a(u, u; \mu)}{\|u\|^2} > 0$$

and the coercivity is uniform with respect to μ in the sense that for some $\bar{\alpha} > 0$ holds $\alpha(\mu) \geq \bar{\alpha}$ for all $\mu \in \mathcal{P}$. Similarly, we assume the parametric linear forms $f(\cdot; \mu), l(\cdot; \mu)$ to be uniformly continuous, i.e. there exist constants $\bar{\gamma}_f, \bar{\gamma}_l < \infty$ such that for all $\mu \in \mathcal{P}$

$$\|l(\cdot; \mu)\|_{X'} \leq \bar{\gamma}_l, \quad \|f(\cdot; \mu)\|_{X'} \leq \bar{\gamma}_f.$$

Example 1 (Possible discontinuity with respect to μ). *Note, that continuity of a , f , and l with respect to u, v does not imply continuity with respect to μ . A simple counterexample can be formulated by $X = \mathbb{R}$, $\mathcal{P} := [0, 2]$, $l : X \times \mathcal{P} \rightarrow \mathbb{R}$ defined as $l(x; \mu) := x \chi_{[1, 2]}(\mu)$ with $\chi_{[1, 2]}$ denoting the indicator function of the specified interval. Obviously, l is continuous with respect to x for all μ , but discontinuous with respect to μ .*

With these assumptions, we can define the full problem that is to be approximated by the subsequent RB-scheme. The full problem can both comprise a continuous PDE in infinite-dimensional function spaces, but as well a finite element discretization of a PDE. The former view is interesting from a theoretical point (how well can solution manifolds in function spaces be approximated), the latter is important from a practical view, as highly resolved discretized PDEs will serve as snapshot suppliers for the reduced basis generation and as the reference solution to compare with.

Definition 2.2 (Full Problem $(P(\mu))$). For $\mu \in \mathcal{P}$ find a solution $u(\mu) \in X$ and output $s(\mu) \in \mathbb{R}$ satisfying

$$\begin{aligned} a(u(\mu), v; \mu) &= f(v; \mu), \quad v \in X, \\ s(\mu) &= l(u(\mu); \mu). \end{aligned}$$

Under the above conditions, one obtains well-posedness and stability of $(P(\mu))$:

Proposition 2.3 (Well-posedness and Stability of $(P(\mu))$). The problem $(P(\mu))$ admits a unique solution satisfying

$$\|u(\mu)\| \leq \frac{\|f(\cdot; \mu)\|_{X'}}{\alpha(\mu)} \leq \frac{\tilde{\gamma}_f}{\tilde{\alpha}}, \quad |s(\mu)| \leq \|l(\cdot; \mu)\|_{X'} \|u(\mu)\| \leq \frac{\tilde{\gamma}_l \tilde{\gamma}_f}{\tilde{\alpha}}.$$

Proof. The existence, uniqueness and bound for $u(\mu)$ follow from the Lax-Milgram theorem, see for instance [8]. Uniform continuity and coercivity then give the parameter-independent bound for $u(\mu)$. The definition of the output functional gives uniqueness for $s(\mu)$ and uniform continuity of l yields the second bound for $s(\mu)$. $\square$

After having ensured solvability, it makes sense to introduce the solution manifold:

Definition 2.4 (Solution Manifold). We introduce the solution manifold $\mathcal{M}$ of the full problem $(P(\mu))$ as

$$\mathcal{M} := \{u(\mu) | u(\mu) \text{ solves } (P(\mu)) \text{ for } \mu \in \mathcal{P}\} \subset X.$$

We use the notion “manifold”, although strictly speaking, it may not be a manifold in the differential geometrical sense, as we do not assume continuity/differentiability of $\mathcal{M}$.

A crucial property for efficient implementation of RB-methods is a parameter-separability of all (bi)linear forms.

Definition 2.5 (Parameter-Separability). We assume the forms a, f, l to be parameter-separable, i.e. there exist coefficient functions $\theta_q^a(\mu) : \mathcal{P} \rightarrow \mathbb{R}$ for $q = 1, \dots, Q_a$ with $Q_a \in \mathbb{N}$ and parameter-independent continuous bilinear forms $a_q(\cdot, \cdot) : X \times X \rightarrow \mathbb{R}$ such that

$$a(u, v; \mu) = \sum_{q=1}^{Q_a} \theta_q^a(\mu) a_q(u, v), \quad \mu \in \mathcal{P}, u, v \in X$$

and similar definitions for l and f with corresponding continuous linear forms l_q, f_q , coefficient functions θ_q^l, θ_q^f and numbers of components Q_l, Q_f .

Note, that in the literature this property is commonly called *affine parameter dependence*. However, this notion is slightly misleading, as the decomposition can be arbitrarily nonlinear (and hence non-affine) with respect to the parameter μ . Therefore, we rather denote it *parameter-separability*.

In the above definition, the constants Q_a, Q_f, Q_l are assumed to be preferably small (e.g. 1–100), as the complexity of the RB-model will explicitly depend on these. We further assume that the coefficient functions $\theta_q^a, \theta_q^f, \theta_q^l$ can be evaluated rapidly.

If such a representation does not exist for a given linear or bilinear form, a parameter separable approximation can be constructed by the Empirical Interpolation Method, we comment on this in Sec. 2.7.

Obviously, boundedness of the coefficient functions θ_a^q and continuity of the components a_q implies uniform continuity of $a(\cdot, \cdot; \mu)$, and similarly for f, l . However, coercivity of components a_q only transfers to coercivity of a under additional assumptions, cf. Exercise 5.3.

Example 2 (Thermal block as instantiation of $(P(\mu))$). *It can easily be verified that the thermal block model satisfies the above assumptions, cf. Exercise 5.4.*

Example 3 ($(P(\mu))$ for matrix equations). *The problem $(P(\mu))$ can also be applied to model order reduction for parametric matrix equations, i.e. solving and reducing systems*

$$\mathbf{A}(\mu)u = \mathbf{b}(\mu)$$

for $\mathbf{A}(\mu) \in \mathbb{R}^{\mathcal{N} \times \mathcal{N}}, \mathbf{b}(\mu) \in \mathbb{R}^{\mathcal{N}}$ for $\mathcal{N} \in \mathbb{N}$. This can simply be obtained by considering $X = \mathbb{R}^{\mathcal{N}}, a(u, v; \mu) := u^T \mathbf{A}(\mu)v$ and $f(v; \mu) := v^T \mathbf{b}(\mu)$ (and ignoring or choosing arbitrary l).

Example 4 ($(P(\mu))$ with given solution). *One can prescribe any arbitrarily complicated parametric function $u : \mathcal{P} \rightarrow X$, which defines a corresponding manifold $\mathcal{M} := \{u(\mu)\}$. Then one can construct an instantiation of $(P(\mu))$ which has this solution. For this we only need to set $a(u, v) := \langle u, v \rangle$ and $f(v; \mu) := \langle u(\mu), v \rangle$ and immediately verify that $u(\mu)$ solves the corresponding $(P(\mu))$. This means that the class $(P(\mu))$ may have arbitrarily complex, nonsmooth or even discontinuous solution manifolds.*

Example 5 ($(P(\mu))$ for $Q_a = 1$). *If $a(\cdot, \cdot; \mu)$ consists of a single component, one can show that $\mathcal{M}$ is contained in an (at most) Q_f -dimensional linear space, cf. Exercise 5.5. Hence, a finite-dimensional RB-space will provide exact approximation.*

We state the following remark on a general misunderstanding of complexity in parametric problems.

Remark 2.6 (Parameter Number and Complexity). *Frequently, problems with many parameters are concluded to have high solution manifold complexity. This is in general wrong: First, as already seen in the parameter redundancy of the thermal block, having many parameters does not necessarily imply a complex solution structure. In the extreme case, one can devise models with arbitrary number of parameters but 1-dimensional solution set, cf. Exercise 5.2. On the other extreme, one can devise models, where 1 parameter induces arbitrary complex solution behavior: In view of Example 4 one can choose an arbitrary connected irregular manifold, and assign a 1-parameter “space-filling-curve” as solution trajectory on this manifold. The misconception that the number of parameters is directly related to the manifold complexity, is very common, even in the model order reduction community. Nevertheless, certainly, in specific examples, the parameter number may influence the solution complexity, such as exemplified for the thermal block later on.*

It is of interest to state properties of $\mathcal{M}$ that may give information on its complexity/approximability. The first insight is that Prop. 2.3 obviously implies the boundedness of the manifold.

It is possible to show a certain regularity of the manifold. First, one can prove Lipschitz-continuity of the manifold, along the lines of [20], cf. Exercise 5.6.

Proposition 2.7 (Lipschitz-Continuity). *If the coefficient functions $\theta_q^a, \theta_q^f, \theta_q^l$ are Lipschitz-continuous with respect to μ , then the forms a, f, l , the solution $u(\mu)$ and $s(\mu)$ are Lipschitz-continuous with respect to μ .*

Further, if the data functions are differentiable, one can even conclude differentiability of the solution manifold by “formally” differentiating $(P(\mu))$. We leave the proof of the following as Exercise 5.7.

Proposition 2.8 (Differentiability, Sensitivity Problem). *If the coefficient functions θ_q^a, θ_q^f are differentiable with respect to μ , then the solution $u : \mathcal{P} \rightarrow X$ is differentiable with respect to μ and the partial derivative (sensitivity derivative) $\partial_{\mu_i} u(\mu) \in X$ for $i = 1, \dots, p$ satisfies the sensitivity problem*

$$a(\partial_{\mu_i} u(\mu), v; \mu) = \tilde{f}_i(v; u(\mu), \mu)$$

for right-hand side $\tilde{f}_i(\cdot; u(\mu), \mu) \in X'$,

$$\tilde{f}_i(\cdot; u(\mu), \mu) := \sum_{q=1}^{Q_f} (\partial_{\mu_i} \theta_q^f(\mu)) f_q(\cdot) - \sum_{q=1}^{Q_a} (\partial_{\mu_i} \theta_q^a(\mu)) a_q(u(\mu), \cdot; \mu).$$

Similar statements hold for higher order differentiability. So, the partial derivatives of $u(\mu)$ satisfy a similar problem as $(P(\mu))$, in particular involving the identical bilinear form, but a right hand side that depends on the lower order derivatives. So we conclude that smoothness of coefficient functions transfers to corresponding smoothness of the solution manifold. Then, the smoother the manifold the better approximability by low-dimensional spaces may be expected.

2.3 Primal RB-Approach

We now formulate two RB-approaches for the above problem class, which have been similarly introduced in [59]. For the moment, we assume to have a low-dimensional space

$$X_N := \text{span}(\Phi_N) = \text{span}\{u(\mu^{(1)}), \dots, u(\mu^{(N)})\} \subset X \quad (1)$$

with a basis $\Phi_N = \{\varphi_1, \dots, \varphi_N\}$ available, which will be called the reduced basis space in the following. The functions $u(\mu^{(i)})$ are suitably chosen *snapshots* of the full problem at parameter samples $\mu^{(i)} \in \mathcal{P}$. We will give details on procedures for their choice in Sec. 2.6. At the moment we only assume that the $\{u(\mu^{(i)})\}_{i=1}^N$ are linearly independent. The first RB-formulation is a straightforward Galerkin projection. It is denoted “primal”, as we will later add a “dual” problem.

Definition 2.9 (Primal RB-problem $(P_N(\mu))$). For $\mu \in \mathcal{P}$ find a solution $u_N(\mu) \in X_N$ and output $s_N(\mu) \in \mathbb{R}$ satisfying

$$\begin{aligned} a(u_N(\mu), v; \mu) &= f(v; \mu), \quad v \in X_N, \\ s_N(\mu) &= l(u_N(\mu)). \end{aligned} \quad (2)$$

Again, well-posedness and stability of the reduced solution follow by Lax-Milgram, even using the same constants as for the full problem, as continuity and coercivity are inherited to subspaces.

Proposition 2.10 (Well-posedness and Stability of $(P_N(\mu))$). The problem $(P_N(\mu))$ admits a unique solution satisfying

$$\|u_N(\mu)\| \leq \frac{\|f(\mu)\|_{X'}}{\alpha(\mu)} \leq \frac{\bar{\gamma}_f}{\bar{\alpha}}, \quad |s_N(\mu)| \leq \|l(\cdot; \mu)\|_{X'} \|u_N(\mu)\| \leq \frac{\bar{\gamma}_l \bar{\gamma}_f}{\bar{\alpha}}.$$

Proof. We verify the applicability of the Lax-Milgram Theorem with the same constants as the full problem using $X_N \subset X$:

$$\begin{aligned} \sup_{u, v \in X_N \setminus \{0\}} \frac{a(u, v; \mu)}{\|u\| \|v\|} &\leq \sup_{u, v \in X \setminus \{0\}} \frac{a(u, v; \mu)}{\|u\| \|v\|} = \gamma(\mu), \\ \inf_{u \in X_N \setminus \{0\}} \frac{a(u, u; \mu)}{\|u\|^2} &\geq \inf_{u \in X \setminus \{0\}} \frac{a(u, u; \mu)}{\|u\|^2} = \alpha(\mu). \end{aligned}$$

Then the argumentation as in Prop. 2.3 applies. $\square$

From a computational view, the problem $(P_N(\mu))$ is solved by a simple linear equation system.

Proposition 2.11 (Discrete RB-Problem). For $\mu \in \mathcal{P}$ and a given reduced basis $\Phi_N = \{\varphi_1, \dots, \varphi_N\}$ define the matrix, right hand side and output vector as

$$\begin{aligned} \mathbf{A}_N(\mu) &:= (a(\varphi_j, \varphi_i; \mu))_{i,j=1}^N \in \mathbb{R}^{N \times N}, \\ \mathbf{f}_N(\mu) &:= (f(\varphi_i; \mu))_{i=1}^N \in \mathbb{R}^N, \quad \mathbf{l}_N(\mu) := (l(\varphi_i; \mu))_{i=1}^N \in \mathbb{R}^N. \end{aligned}$$

Solve the following linear system for $\mathbf{u}_N(\mu) = (u_{N,i})_{i=1}^N \in \mathbb{R}^N$:

$$\mathbf{A}_N(\mu) \mathbf{u}_N(\mu) = \mathbf{f}_N(\mu).$$

Then, the solution of $(P_N(\mu))$ is obtained by

$$u_N(\mu) = \sum_{j=1}^N u_{N,j} \varphi_j, \quad s_N(\mu) = \mathbf{l}_N^T(\mu) \mathbf{u}_N(\mu). \quad (3)$$

Proof. Using linearity, it directly follows that $u_N(\mu)$ and $s_N(\mu)$ from (3) satisfy $(P_N(\mu))$. $\square$

Interestingly, in addition to analytical stability from Prop. 2.10 we can also guarantee algebraic stability by using an orthonormal reduced basis. This means we do not have to use snapshots directly as reduced basis vectors, but Φ_N can be a post-processed set of snapshots, as long as (1) holds. For example, a standard Gram-Schmidt orthonormalization can be performed to obtain an orthonormal reduced basis Φ_N .

Proposition 2.12 (Algebraic Stability for Orthonormal Basis). *If $a(\cdot, \cdot; \mu)$ is symmetric and Φ_N is orthonormal, then the condition number of $\mathbf{A}_N(\mu)$ is bounded (independently of N) by*

$$\text{cond}_2(\mathbf{A}_N(\mu)) = \|\mathbf{A}_N(\mu)\| \|\mathbf{A}_N(\mu)^{-1}\| \leq \frac{\gamma(\mu)}{\alpha(\mu)}.$$

Proof. As $\mathbf{A}_N$ is symmetric and positive definite we have $\text{cond}_2(\mathbf{A}_N) = \lambda_{\max}/\lambda_{\min}$ with largest/smallest magnitude eigenvalue of $\mathbf{A}_N$. Let $\mathbf{u} = (u_i)_{i=1}^N \in \mathbb{R}^N$ be an eigenvector of $\mathbf{A}_N$ for eigenvalue $\lambda_{\max}$ and set $u := \sum_{i=1}^N u_i \varphi_i \in X$. Then due to orthonormality we obtain

$$\|u\|^2 = \left\langle \sum_{i=1}^N u_i \varphi_i, \sum_{j=1}^N u_j \varphi_j \right\rangle = \sum_{i,j=1}^N u_i u_j \langle \varphi_i, \varphi_j \rangle = \sum_{i=1}^N u_i^2 = \|\mathbf{u}\|^2,$$

by definition of $\mathbf{A}_N$ and continuity we get

$$\lambda_{\max} \|\mathbf{u}\|^2 = \mathbf{u}^T \mathbf{A}_N \mathbf{u} = a \left(\sum_{i=1}^N u_i \varphi_i, \sum_{j=1}^N u_j \varphi_j \right) = a(u, u) \leq \gamma \|u\|^2$$

and we conclude that $\lambda_{\max} \leq \gamma$. Similarly, one can show that $\lambda_{\min} \geq \alpha$, which then gives the desired statement. $\square$

This uniform stability bound is a relevant advantage over a non-orthonormal snapshot basis. In particular one can easily realize that snapshots of “close” parameters will result in almost colinear snapshots leading to similar columns and therefore ill-conditioning of the reduced system matrix. But still, for small reduced bases (e.g. 1–10), the orthonormalization can be omitted in order to prevent additional numerical errors of the orthonormalization procedure.

Remark 2.13 (Difference of FEM to RB). *At this point we can note a few distinct differences between the reduced problem and a discretized full problem. For this denote $\mathbf{A} \in \mathbb{R}^{\mathcal{N} \times \mathcal{N}}$ for some large $\mathcal{N} \in \mathbb{N}$ the Finite Element (or Finite Volume, Discontinuous Galerkin, etc.) matrix of the linear system for $(P(\mu))$. Then,*

- the RB-matrix $\mathbf{A}_N \in \mathbb{R}^{N \times N}$ is small, but typically dense in contrast to $\mathbf{A}$ which is large but typically sparse,
- the condition of the matrix $\mathbf{A}_N$ does not deteriorate with growing N if an orthonormal basis is used, in contrast to the high-dimensional $\mathbf{A}$, whose condition number typically grows polynomially in $\mathcal{N}$.

The RB-approximation will always be as good as the best-approximation, up to a constant. This is a simple version of the Lemma of Céa.

Proposition 2.14 (Céa, Relation to Best-Approximation). *For all $\mu \in \mathcal{P}$ holds*

$$\|u(\mu) - u_N(\mu)\| \leq \frac{\gamma(\mu)}{\alpha(\mu)} \inf_{v \in X_N} \|u(\mu) - v\|. \quad (4)$$

If additionally $a(\cdot, \cdot; \mu)$ is symmetric, we have the sharpened bound

$$\|u(\mu) - u_N(\mu)\| \leq \sqrt{\frac{\gamma(\mu)}{\alpha(\mu)}} \inf_{v \in X_N} \|u(\mu) - v\|. \quad (5)$$

Proof. For all $v \in X_N$ continuity and coercivity result in

$$\begin{aligned} \alpha \|u - u_N\|^2 &\leq a(u - u_N, u - u_N) = a(u - u_N, u - v) + a(u - u_N, v - u_N) \\ &= a(u - u_N, u - v) \leq \gamma \|u - u_N\| \|u - v\|, \end{aligned}$$

where we used Galerkin orthogonality $a(u - u_N, v - u_N) = 0$ which follows from $(P(\mu))$ and $(P_N(\mu))$ as $v - u_N \in X_N$. For the sharpened bound 5 we refer to [57] or Exercise 5.8. $\square$

Similar best-approximation statements are known for interpolation techniques, but the corresponding constants mostly diverge to ∞ as the dimension of the approximating space N grows. For the RB-approximation, the constant does not grow with N . This is the conceptional advantage of RB-approximation by Galerkin projection rather than some other interpolation techniques.

For error analysis the following error-residual relation is important, which states that the error satisfies a variational problem with the same bilinear form, but the residual as right hand side.

Proposition 2.15 (Error-Residual Relation). *For $\mu \in \mathcal{P}$ we define the residual $r(\cdot; \mu) \in X'$ via*

$$r(v; \mu) := f(v; \mu) - a(u_N(\mu), v; \mu), \quad v \in X. \quad (6)$$

Then, the error $e(\mu) := u(\mu) - u_N(\mu) \in X$ satisfies

$$a(e, v; \mu) = r(v; \mu), \quad v \in X. \quad (7)$$

Proof. $a(e, v; \mu) = a(u, v; \mu) - a(u_N, v; \mu) = f(v) - a(u_N, v; \mu) = r(v; \mu)$. $\square$

Hence, the residual in particular vanishes on X_N as $X_N \subset \ker(r(\cdot; \mu))$.

A basic consistency property for an RB-scheme is reproduction of solutions. The following statement follows trivially from the error estimators which will be introduced soon. But in case of absence of error estimators for an RB-scheme, this reproduction property still can be investigated. It states that if a full solution happens to be in the reduced space, then the RB-scheme will identify this full solution as the reduced solution, hence give zero error.

Proposition 2.16 (Reproduction of Solutions). *If $u(\mu) \in X_N$ for some $\mu \in \mathcal{P}$, then $u_N(\mu) = u(\mu)$.*

Proof. If $u(\mu) \in X_N$, then $e = u(\mu) - u_N(\mu) \in X_N$ and we obtain by coercivity and $(P(\mu))$ and $(P_N(\mu))$

$$\alpha(\mu) \|e\|^2 \leq a(e, e; \mu) = a(u(\mu), e; \mu) - a(u_N(\mu), e; \mu) = f(e; \mu) - f(e; \mu) = 0,$$

hence $e = 0$. □

This is a trivial, but useful statement for at least two reasons which we state as remarks.

Remark 2.17 (Validation of RB-scheme). *On the practical side, the reproduction property is useful to validate the implementation of an RB-scheme: Choose Φ_N directly as snapshot basis, i.e. $\varphi_i = u(\mu^{(i)})$ without orthonormalization and set $\mu = \mu^{(i)}$, then the RB-scheme must return $\mathbf{u}_N(\mu) = \mathbf{e}_i$, the i -th unit vector, as $u_N(\mu) = \sum_{n=1}^N \delta_{ni} \varphi_n$ (with δ_{ni} denoting the Kronecker δ) obviously is the solution expansion.*

Remark 2.18 (Uniform convergence of RB-approximation). *From a theoretical view-point we can conclude convergence of the RB-approximation to the full continuous problem: We see that the RB-solution $u_N : \mathcal{P} \rightarrow X$ is interpolating the manifold $\mathcal{M}$ at the snapshot parameters $\mu^{(i)}$. Assume that $\mathcal{P}$ is compact and snapshot parameter samples are chosen, such that the sets $S_N := \{\mu^{(1)}, \dots, \mu^{(N)}\} \subset \mathcal{P}$ get dense in $\mathcal{P}$ for $N \rightarrow \infty$, i.e. the so called fill-distance h_N tends to zero:*

$$h_N := \sup_{\mu \in \mathcal{P}} \text{dist}(\mu, S_N), \quad \lim_{N \rightarrow \infty} h_N = 0.$$

Here, $\text{dist}(\mu, S_N) := \min_{\mu' \in S_N} \|\mu - \mu'\|$ denotes the distance of the point μ from the set S_N . If the data functions are Lipschitz continuous, one can show as in Prop. 2.7 that $u_N : \mathcal{P} \rightarrow X_N$ is Lipschitz continuous with Lipschitz-constant L_u independent of N . Then, obviously, for all N, μ and “closest” $\mu^* = \arg \min_{\mu' \in S_N} \|\mu - \mu'\|$

$$\begin{aligned} \|u(\mu) - u_N(\mu)\| &\leq \|u(\mu) - u(\mu^*)\| + \|u(\mu^*) - u_N(\mu^*)\| + \|u_N(\mu) - u_N(\mu^*)\| \\ &\leq L_u \|\mu - \mu^*\| + 0 + L_u \|\mu - \mu^*\| \leq 2h_N L_u. \end{aligned}$$

Therefore, we obtain uniform convergence

$$\lim_{N \rightarrow \infty} \sup_{\mu \in \mathcal{P}} \|u(\mu) - u_N(\mu)\| = 0.$$

Note, however, that this convergence rate is linear in h_N and thus is of no practical value, as h_N decays much too slow with N , and N must be very large to guarantee a small error. In Sec. 2.6 we will see that a more clever choice of $\mu^{(i)}$ can even result in exponential convergence.

We now turn to an important topic in RB-methods, namely the *certification* by a-posteriori error control. This is also based on the residual. We assume to have

a rapidly computable lower bound $\alpha_{\text{LB}}(\mu)$ for the coercivity constant available and that α_{LB} is still large in the sense that it is bounded away from zero

$$0 < \bar{\alpha} \leq \alpha_{\text{LB}}(\mu).$$

This can be assumed without loss of generality, as in case of $\alpha_{\text{LB}}(\mu) < \bar{\alpha}$ we should better choose $\alpha_{\text{LB}}(\mu) = \bar{\alpha}$ and obtain a larger lower bound constant (assuming $\bar{\alpha}$ to be computable).

Proposition 2.19 (A-posteriori Error Bounds). *Let $\alpha_{\text{LB}}(\mu) > 0$ be a computable lower bound for $\alpha(\mu)$. Then we have for all $\mu \in \mathcal{P}$*

$$\|u(\mu) - u_N(\mu)\| \leq \Delta_u(\mu) := \frac{\|r(\cdot; \mu)\|_{X'}}{\alpha_{\text{LB}}(\mu)}, \quad (8)$$

$$|s(\mu) - s_N(\mu)| \leq \Delta_s(\mu) := \|l(\cdot; \mu)\|_{X'} \Delta_u(\mu). \quad (9)$$

Proof. The case $e = 0$ is trivial, hence we assume nonzero error. Testing the error-residual equation with e yields

$$\alpha(\mu) \|e\|^2 \leq a(e, e; \mu) = r(e; \mu) \leq \|r(\cdot; \mu)\|_{X'} \|e\|.$$

Division by $\|e\|$ and α yields the bound for $\|e\|$. The bound for the output error follows by continuity from

$$|s(\mu) - s_N(\mu)| = |l(u(\mu); \mu) - l(u_N(\mu); \mu)| \leq \|l(\cdot; \mu)\|_{X'} \|u(\mu) - u_N(\mu)\|,$$

which concludes the proof. $\square$

Note, that bounding the error by the residual is a well known technique in FEM analysis for comparing a FEM-solution to the analytical solution. However, in that case X is infinite-dimensional, and the norm $\|r\|_{X'}$ is not available analytically. In our case, by using X to be a fine discrete FEM space, the residual norm becomes a computable quantity, which can be computed *after* the reduced solution $u_N(\mu)$ is available, hence it is an *a-posteriori* bound.

The above technique is an example of a general procedure for obtaining error bounds for RB-methods: Show that the RB-error satisfies a problem similar to the original problem, but with a residual as inhomogeneity. Then apply an *a-priori* stability estimate, to get an error bound in terms of a residual norm, which is computable in the RB-setting.

As the bounds are provable upper bounds to the error, they are denoted *rigorous* error bounds. The availability of a-posteriori error bounds is the motivation to denote the approach a *certified* RB-method, as we not only obtain an RB-approximation but simultaneously a certification by a guaranteed error bound.

Having a bound, the question arises, how tight this bound is. A first desirable property of an error bound is that it should be zero if the error is zero, hence we can a-posteriori identify exact approximation.

Corollary 2.20 (Vanishing Error Bound). *If $u(\mu) = u_N(\mu)$ then $\Delta_u(\mu) = \Delta_s(\mu) = 0$.*

Proof. As $0 = a(0, v; \mu) = a(e, v) = r(v; \mu)$, we see that

$$\|r(\cdot; \mu)\|_{X'} := \sup_u \|r(u; \mu)\| / \|u\| = 0.$$

This implies that $\Delta_u(\mu) = 0$ and $\Delta_s(\mu) = 0$. $\square$

This may give hope that the quotient of error bounds and true error behaves well. In particular, the factor of overestimation can be investigated, and, ideally, be bounded by a small constant. The error bounds are then called *effective*. This is possible for $\Delta_u(\mu)$ in our scenario and the so called effectivity can be bounded by the continuity and coercivity constant. Thanks to the uniform continuity and coercivity, this is even parameter-independent.

Proposition 2.21 (Effectivity Bound). *The effectivity $\eta_u(\mu)$ is defined and bounded by*

$$\eta_u(\mu) := \frac{\Delta_u(\mu)}{\|u(\mu) - u_N(\mu)\|} \leq \frac{\gamma(\mu)}{\alpha_{\text{LB}}(\mu)} \leq \frac{\bar{\gamma}}{\bar{\alpha}}. \quad (10)$$

Proof. Let $v_r \in X$ denote the Riesz-representative of $r(\cdot; \mu)$, i.e. we have

$$\langle v_r, v \rangle = r(v; \mu), \quad v \in X, \quad \|v_r\| = \|r(\cdot; \mu)\|_{X'}.$$

Then, we obtain via the error-residual-equation (7) and continuity

$$\|v_r\|^2 = \langle v_r, v_r \rangle = r(v_r; \mu) = a(e, v_r; \mu) \leq \gamma(\mu) \|e\| \|v_r\|.$$

Hence $\frac{\|v_r\|}{\|e\|} \leq \gamma(\mu)$. We then conclude

$$\eta_u(\mu) = \frac{\Delta_u(\mu)}{\|e\|} = \frac{\|v_r\|}{\alpha_{\text{LB}}(\mu) \|e\|} \leq \frac{\gamma(\mu)}{\alpha_{\text{LB}}(\mu)}$$

and obtain the parameter-independent bound via uniform continuity and coercivity. $\square$

Note, that in view of this statement, Cor. 2.20 is a trivial corollary. Still, the property stated in Cor. 2.20 has a value of its own, and in more complex RB-scenarios without effectivity bounds it may be all one can get. Due to the proven reliability and effectivity of the error bounds, these are also denoted *error estimators*, as they obviously are equivalent to the error up to suitable constants.

In addition to absolute error bounds, it is also possible to derive relative error and effectivity bounds. We again refer to [57] for corresponding proofs and similar statements for other error measures. See also [27] or Exercise 5.9.

Proposition 2.22 (Relative Error Bound and Effectivity). *We have for all $\mu \in \mathcal{P}$*

$$\begin{aligned} \frac{\|u(\mu) - u_N(\mu)\|}{\|u(\mu)\|} &\leq \Delta_u^{\text{rel}}(\mu) := 2 \cdot \frac{\|r(\cdot; \mu)\|_{X'}}{\alpha_{\text{LB}}(\mu)} \cdot \frac{1}{\|u_N(\mu)\|}, \\ \eta_u^{\text{rel}}(\mu) &:= \frac{\Delta_u^{\text{rel}}}{\|e(\mu)\|/\|u(\mu)\|} \leq 3 \cdot \frac{\gamma(\mu)}{\alpha_{\text{LB}}(\mu)}, \end{aligned} \quad (11)$$

if $\Delta_u^{\text{rel}}(\mu) \leq 1$.

Hence, these relative bounds are valid, if the error estimator is sufficiently small.

One can put on different “glasses” when analyzing an error, i.e. use different norms, and perhaps obtain sharper bounds. This is possible for the current case by using the (parameter-dependent) energy norm. For this we assume that a is symmetric and define

$$\langle u, v \rangle_\mu := a(u, v; \mu).$$

This form is positive definite by coercivity of a . Hence, $\langle \cdot, \cdot \rangle_\mu$ is a scalar product and induces the *energy norm*

$$\|u\|_\mu := \sqrt{\langle u, u \rangle_\mu}.$$

By coercivity and continuity of a one can easily see that the energy norm is equivalent to the norm on X by

$$\sqrt{\alpha(\mu)}\|u\| \leq \|u\|_\mu \leq \sqrt{\gamma(\mu)}\|u\|, \quad u \in X. \quad (12)$$

With respect to this norm, one can derive an improved error bound and effectivity. We omit the proof, and refer to [57] or [27] and Exercise 5.10.

Proposition 2.23 (Energy-norm Error Bound and Effectivity). *For $\mu \in \mathcal{P}$ with symmetric $a(\cdot, \cdot; \mu)$ we have*

$$\begin{aligned} \|u(\mu) - u_N(\mu)\|_\mu &\leq \Delta_u^{\text{en}}(\mu) := \frac{\|r(\cdot; \mu)\|_{X'}}{\sqrt{\alpha_{\text{LB}}(\mu)}}, \\ \eta_u^{\text{en}}(\mu) &:= \frac{\Delta_u^{\text{en}}}{\|e\|_\mu} \leq \sqrt{\frac{\gamma(\mu)}{\alpha_{\text{LB}}(\mu)}}. \end{aligned} \quad (13)$$

As $\gamma(\mu)/\alpha_{\text{LB}}(\mu) \geq 1$ this is an improvement by a square root compared to (10).

The energy norm allows another improvement in the RB-methodology: By choosing a specific $\bar{\mu} \in \mathcal{P}$ one can choose $\|\cdot\| := \|\cdot\|_{\bar{\mu}}$ as norm on X . Then by definition, one obtains $\gamma(\bar{\mu}) = 1 = \alpha(\bar{\mu})$. This means that for the selected parameter the effectivity is $\eta_u(\bar{\mu}) = 1$, hence the error bound exactly corresponds to the error norm. In this sense, the error bound is optimal. Assuming continuity of $\alpha(\mu), \gamma(\mu)$, one can therefore expect that choosing this norm on X will give highly effective RB-error bounds also in an environment of $\bar{\mu}$.

We continue with further specialization. For the special case of a *compliant* problem, the above RB-scheme ($P_N(\mu)$) turns out to be very good; we obtain effectivities and an output bound that is quadratic in $\Delta_u(\mu)$ instead of only linear. The proof of this statement can be found in [57] or follows as a special instance from Prop. 2.27 in the next section, cf. Remark 2.28. Still, as this bound is somehow a central statement, we give the proof.

Proposition 2.24 (Output Error Bound and Effectivity For “Compliant” Case). *If $a(\cdot, \cdot; \mu)$ is symmetric and $l = f$ (the so called “compliant” case), we obtain the improved output bound*

$$0 \leq s(\mu) - s_N(\mu) \leq \Delta_s'(\mu) := \frac{\|r(\cdot; \mu)\|_{X'}^2}{\alpha_{\text{LB}}(\mu)} = \alpha_{\text{LB}}(\mu) \Delta_u(\mu)^2 \quad (14)$$

and effectivity bound

$$\eta'_s(\mu) := \frac{\Delta'_s(\mu)}{s(\mu) - s_N(\mu)} \leq \frac{\gamma(\mu)}{\alpha_{\text{LB}}(\mu)} \leq \frac{\bar{\gamma}}{\bar{\alpha}}. \quad (15)$$

Proof. Using $a(u_N, e) = 0$ due to Galerkin orthogonality we obtain (omitting μ for brevity)

$$s - s_N = l(u) - l(u_N) = l(e) = f(e) \quad (16)$$

$$= f(e) - a(u_N, e) = r(e) = a(e, e). \quad (17)$$

Coercivity then implies the first inequality of (14). The second inequality and the last equality of (14) follow from the error-residual relation and the bound for u :

$$a(e, e) = r(e) \leq \|r\| \|e\| \leq \|r\| \Delta_u = \|r\| \frac{\|r\|}{\alpha_{\text{LB}}} = \alpha_{\text{LB}} \Delta_u^2. \quad (18)$$

For the effectivity bound (15) we first note that with Cauchy-Schwarz and norm-equivalence (12) the Riesz-representative v_r satisfies

$$\|v_r\|^2 = \langle v_r, v_r \rangle = r(v_r) = a(e, v_r) = \langle e, v_r \rangle_\mu \leq \|e\|_\mu \|v_r\|_\mu \leq \|e\|_\mu \sqrt{\gamma} \|v_r\|.$$

Assuming $v_r \neq 0$ division by $\|v_r\|$ yields

$$\|r\|_{X'} = \|v_r\| \leq \|e\|_\mu \sqrt{\gamma}.$$

For $v_r = 0$ this inequality is trivially satisfied. This allows to conclude using the definitions and (17)

$$\eta'_s(\mu) = \frac{\Delta_s}{s - s_N} = \frac{\|r\|_{X'}^2 / \alpha}{a(e, e)} = \frac{\|r\|_{X'}^2}{\alpha \|e\|_\mu^2} \leq \frac{\gamma \|e\|_\mu^2}{\alpha \|e\|_\mu^2} \leq \frac{\bar{\gamma}}{\bar{\alpha}}.$$

□

Note, that the proposition gives a definite sign on the output error, i.e. we always have $s_N(\mu) \leq s(\mu)$.

We will conclude this section with some experimental results which illustrate the theoretical findings. These results can be reproduced via the package *RBmatlab* which is available for download at www.morepas.org. It is a package providing different grid types for spatial discretization, different discretization schemes for PDEs, and various models and implementations of RB-schemes. One example is the thermal block model, which is also realized in that package, in particular the plots in Fig. 3 and the subsequent experiments can be reproduced by the program `rb_tutorial.m`. In the following we recommend to inspect the source code of that program and verify the following results by running different parts of the script. If the reader does not want to install the complete package, a standalone

script `rb_tutorial_standalone.m` using some precomputed data files offers the same functionality. These files are also accessible via www.morepas.org.

We consider the thermal block with $B_1 = B_2 = 2$, $\mu_{\min} = 1/\mu_{\max} = 0.1$, choose 5 sampling points $\mu^{(j)} = (0.1 + 0.5(j-1), c, c, c)^T$, $j = 1, \dots, 5$ with $c = 0.1$ and plot the error estimator $\Delta_u(\mu)$ and true error $\|u(\mu) - u_N(\mu)\|$ for $\mu = (\mu_1, c, c, c)$ over μ_1 . The results are depicted in Fig. 4a). We can see that the error estimator is finely re-

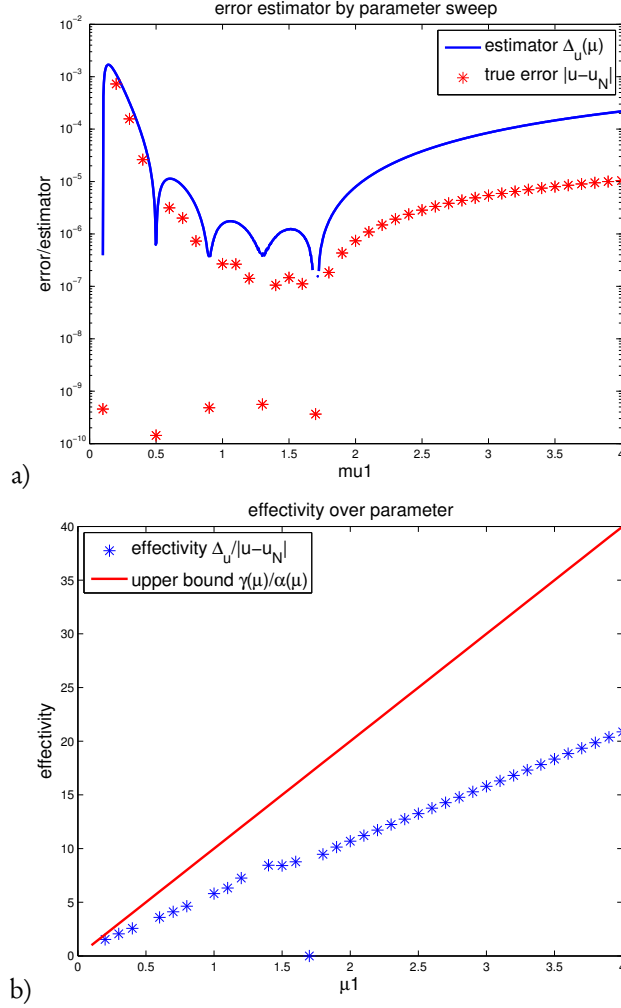


Figure 4: Illustration of a) error and error bound and b) effectivity and effectivity bound over parameter.

solved, as a parameter sweep is computationally cheap thanks to the reduced model. The true error has been sampled more coarsely, as solving the full problem is more

tedious. We see that the error bound indeed is an upper bound for the error, confirming the rigor. Further, we see that the true error indeed is (numerically) zero for the chosen sampling points, due to the reproduction of solutions, Prop. 2.16. Also the error bound is zero in these points, as is expected by the vanishing error bound property, Cor. 2.20. Finally, the error between sampling points is growing for low-value intervals. This reflects the requirement that for small diffusivity coefficients denser sampling is necessary for uniform error distribution. This fact will be supported later by some a-priori analysis.

If we look at the effectivities in Fig. 4b), we indeed see that $\eta_\mu(\mu)$ is only well defined for parameters with nonzero error and it is bounded from above by $\gamma(\mu)/\alpha(\mu)$ in accordance with Prop. 2.21. The values of the effectivities are only in the order of 10, which is considered as quite good.

2.4 Primal-Dual RB-Approach

As seen in the previous section, the output error bound $\Delta_s(\mu)$ is scaling linearly with Δ_μ for the general case, Prop. 2.19, and quadratically for the compliant case, Prop. 2.24. By involving a goal-oriented strategy [3] via a corresponding dual problem, one can improve the output and output error estimation in the sense that the output error bounds will show this “quadratic” behavior also for non-compliant problems. For an early reference on the presented RB-approach we refer to [59]. We first define the full dual problem.

Definition 2.25 (Full Dual Problem ($P'(\mu)$)). *For $\mu \in \mathcal{P}$ find a solution $u^{\text{du}}(\mu) \in X$ of*

$$a(v, u^{\text{du}}(\mu); \mu) = -l(v; \mu), \quad v \in X.$$

Again, well-posedness and stability are guaranteed due to coercivity and continuity.

We assume that an RB-space $X_N^{\text{du}} \subset X$ with dimension N^{du} , not necessarily equal to N , is available, cf. Sec. 2.6 for comments on corresponding basis-generation strategies. Then we can define the primal-dual RB-approach.

Definition 2.26 (Primal-Dual RB-problem ($P'_N(\mu)$)). *For $\mu \in \mathcal{P}$ let $u_N(\mu) \in X_N$ be the solution of ($P_N(\mu)$). Then, find a solution $u_N^{\text{du}}(\mu) \in X_N^{\text{du}}$ and output $s'_N(\mu) \in \mathbb{R}$ satisfying*

$$\begin{aligned} a(v, u_N^{\text{du}}(\mu); \mu) &= -l(v; \mu), \quad v \in X_N^{\text{du}}, \\ s'_N(\mu) &= l(u_N(\mu)) - r(u_N^{\text{du}}; \mu). \end{aligned}$$

Again, well-posedness and stability follow by coercivity and continuity.

We observe that compared to the primal output $s_N(\mu)$ from (2) we have an output estimate $s'_N(\mu)$ using a “correction term” given by the primal residual evaluated at the dual solution. This “correction” allows to derive sharper error bounds.

For the primal solution $u_N(\mu)$ the error bound (8) and effectivity (10) still are valid. For the dual variable and the corrected output we obtain the following.

Proposition 2.27 (A-posteriori Error and Effectivity Bounds). *For all $\mu \in \mathcal{P}$ we introduce the dual residual*

$$r^{\text{du}}(v; \mu) := -l(v; \mu) - a(v, u_N^{\text{du}}(\mu); \mu), \quad v \in X$$

and obtain a-posteriori error bounds

$$\begin{aligned} \|u^{\text{du}}(\mu) - u_N^{\text{du}}(\mu)\| &\leq \Delta_u^{\text{du}}(\mu) := \frac{\|r^{\text{du}}(\cdot; \mu)\|_{X'}}{\alpha_{\text{LB}}(\mu)}, \\ |s(\mu) - s'_N(\mu)| &\leq \Delta'_s(\mu) := \frac{\|r^{\text{du}}(\cdot; \mu)\|_{X'} \|r(\cdot; \mu)\|_{X'}}{\alpha_{\text{LB}}(\mu)} \end{aligned} \quad (19)$$

$$= \alpha_{\text{LB}}(\mu) \Delta_u(\mu) \Delta_u^{\text{du}}(\mu), \quad (20)$$

and the effectivity bound

$$\eta_u^{\text{du}}(\mu) := \frac{\Delta_u^{\text{du}}(\mu)}{\|u^{\text{du}}(\mu) - u_N^{\text{du}}(\mu)\|} \leq \frac{\gamma(\mu)}{\alpha_{\text{LB}}(\mu)} \leq \frac{\tilde{\gamma}}{\tilde{\alpha}}.$$

Proof. The bound and effectivity for the dual solution $u_N^{\text{du}}(\mu)$ follow with identical arguments as for the primal solution. For the improved output bound we first note

$$\begin{aligned} s - s'_N &= l(u) - l(u_N) + r(u_N^{\text{du}}) = l(u - u_N) + r(u_N^{\text{du}}) \\ &= -a(u - u_N, u_N^{\text{du}}) + f(u_N^{\text{du}}) - a(u_N, u_N^{\text{du}}) \\ &= -a(u - u_N, u_N^{\text{du}}) + a(u, u_N^{\text{du}}) - a(u_N, u_N^{\text{du}}) \\ &= -a(u - u_N, u_N^{\text{du}} - u_N^{\text{du}}). \end{aligned} \quad (21)$$

Therefore, with the dual error $e^{\text{du}} := u^{\text{du}} - u_N^{\text{du}}$ we obtain

$$\begin{aligned} |s - s'_N| &\leq |a(e, e^{\text{du}})| = |r(e^{\text{du}})| \leq \|r\|_{X'} \|e^{\text{du}}\| \\ &\leq \|r\|_{X'} \cdot \Delta_u^{\text{du}} \leq \|r\|_{X'} \|r^{\text{du}}\|_{X'} / \alpha_{\text{LB}} = \alpha_{\text{LB}} \Delta_u \Delta_u^{\text{du}}. \end{aligned}$$

□

Assuming $\Delta_u(\mu) \approx \Delta_u^{\text{du}}(\mu) \approx h \ll 1$ we see a quadratic dependence of Δ'_s on h in contrast to the simple linear dependence in the output estimate and bound of Prop. 2.19. Hence, the primal-dual approach is expected to give much better output error bounds.

Example 6 (Missing Effectivity Bounds for $\Delta_s(\mu), \Delta'_s(\mu)$). *Note, that in the general non-compliant case, we cannot hope to obtain effectivity bounds for the output error estimators. The reason for this is that $s(\mu) - s_N(\mu)$ or $s(\mu) - s'_N(\mu)$ may be zero, while*

$\Delta_s(\mu), \Delta'_s(\mu)$ are not. Hence, the effectivity as quotient of these quantities is not well defined: Choose vectors $v_l \perp v_f \in X$ and a subspace $X_N^{\text{du}} = X_N \perp \{v_l, v_f\}$. For $a(u, v) := \langle u, v \rangle, f(v) := \langle v_f, v \rangle$ and $l(v) := -\langle v_l, v \rangle$ we obtain $u = v_f$ as primal, $u^{\text{du}} = v_l$ as dual solution and $u_N = 0, u_N^{\text{du}} = 0$ as RB-solutions. Hence $e = v_f, e^{\text{du}} = v_l$. This yields $s = s_N = 0$, but $r \neq 0$ and $r^{\text{du}} \neq 0$ hence $\Delta_s(\mu) > 0$. Similarly, from (21) we obtain $s - s'_N = -a(e, e^{\text{du}}) = \langle v_f, v_l \rangle = 0$. But $r \neq 0$ and $r^{\text{du}} \neq 0$, hence $\Delta'_s(\mu) > 0$. So, further assumptions, such as in the compliant case, are required in order to derive output error bound effectivities.

Remark 2.28 (Equivalence of $(P_N(\mu))$ and $(P'_N(\mu))$ for Compliant Case). In Prop. 2.24 we have given a quadratic output bound statement for the compliant case. In fact, that bound is a simple corollary of Prop. 2.27: We first note that $(P_N(\mu))$ and $(P'_N(\mu))$ are equivalent for the compliant case and assuming $X_N = X_N^{\text{du}}$: As $f = l$ and a is symmetric, $u_N^{\text{du}}(\mu) = -u_N(\mu)$ solves the dual problem and $\Delta_u(\mu) = \Delta_u^{\text{du}}(\mu)$. The residual correction term vanishes, $r(u_N^{\text{du}}) = 0$ as $X_N \in \ker(r)$ and therefore $s'_N(\mu) = s_N(\mu)$. Similarly, equivalence of $(P(\mu))$ and $(P'(\mu))$ holds, hence $e^{\text{du}} = -e$. Then, from (21) we conclude that $s - s_N = -a(e, e^{\text{du}}) = a(e, e) \geq 0$ and the second inequality in (14) follows from (20). Therefore, $(P_N(\mu))$ is fully sufficient in these cases, and the additional technical burden of the primal-dual-approach can be circumvented.

2.5 Offline/Online Decomposition

We now address computational aspects of the RB-methodology. We will restrict ourselves to the primal RB-problem, the primal-dual approach can be treated similarly. As the computational procedure will assume that $(P(\mu))$ is a high-dimensional discrete problem, we will first introduce the corresponding notation. We assume that $X = \text{span}(\psi_i)_{i=1}^{\mathcal{N}}$ is spanned by a large number of basis functions ψ_i . We introduce the system matrix, inner product matrix, and functional vectors as

$$\mathbf{A}(\mu) := (a(\psi_j, \psi_i; \mu))_{i,j=1}^{\mathcal{N}} \in \mathbb{R}^{\mathcal{N} \times \mathcal{N}}, \quad \mathbf{K} := (\langle \psi_i, \psi_j \rangle)_{i,j=1}^{\mathcal{N}} \in \mathbb{R}^{\mathcal{N} \times \mathcal{N}}, \quad (22)$$

$$\mathbf{f}(\mu) := (f(\psi_i; \mu))_{i=1}^{\mathcal{N}} \in \mathbb{R}^{\mathcal{N}}, \quad \mathbf{l}(\mu) := (l(\psi_i; \mu))_{i=1}^{\mathcal{N}} \in \mathbb{R}^{\mathcal{N}}. \quad (23)$$

Then, the full problem $(P(\mu))$ can be solved by determining the coefficient vector $\mathbf{u} = (u_i)_{i=1}^{\mathcal{N}} \in \mathbb{R}^{\mathcal{N}}$ for $u(\mu) = \sum_{j=1}^{\mathcal{N}} u_j \psi_j$ and output from

$$\mathbf{A}(\mu)\mathbf{u}(\mu) = \mathbf{f}(\mu), \quad s(\mu) = \mathbf{l}(\mu)^T \mathbf{u}(\mu). \quad (24)$$

We do not further limit the type of discretization. The system matrix may be obtained from a Finite Element, Finite Volume or Discontinuous Galerkin discretization. Typically, $\mathbf{A}(\mu)$ is a sparse matrix, which is always obtained if local differential operators of the PDE are discretized with basis functions of local support. However, the RB-methodology can in principle also be applied to discretizations resulting in full system matrices, e.g. integral equations or equations with non-local differential terms.

Let us first start with a rough complexity consideration for computing a full and reduced solution. We assume that a single solution of $(P(\mu))$ via (24) requires $\mathcal{O}(\mathcal{N}^2)$ operations (e.g. resulting from $\mathcal{N}$ steps of an iterative solver based on $\mathcal{O}(\mathcal{N})$ for each sparse matrix-vector multiplication). In contrast, the dense reduced problem in Prop. 2.11 is solvable in $\mathcal{O}(N^3)$ (assuming direct inversion of $\mathbf{A}_N$ or N steps of an iterative solver based on $\mathcal{O}(N^2)$ for each matrix-vector multiplication). Hence, we clearly see that the RB-approach requires $N \ll \mathcal{N}$ to realize a computational advantage.

Let us collect the relevant steps for the computation of an RB-solution, (and not consider orthonormalization or the error estimators for the moment):

1. N snapshot computations via $(P(\mu^i))$: $\mathcal{O}(N\mathcal{N}^2)$
2. N^2 evaluations of $a(\varphi_j, \varphi_i; \mu)$: $\mathcal{O}(N^2\mathcal{N})$
3. N evaluations of $f(\varphi_i; \mu)$: $\mathcal{O}(N\mathcal{N})$
4. Solution of the $N \times N$ system $(P_N(\mu))$: $\mathcal{O}(N^3)$.

So, RB-procedures clearly *do not pay off* if a solution for a single parameter μ is required. But in case of multiple solution queries, the RB-approach will pay off due to a so called offline/online decomposition, as already mentioned in the introduction. During the *offline phase* μ -independent, high-dimensional quantities are precomputed. The operation count typically depends on $\mathcal{N}$, hence this phase is *expensive*, but is only performed *once*. During the *online phase* which is performed for *many* parameters $\mu \in \mathcal{P}$, the offline data is combined to give the small μ -dependent discretized reduced system, and the reduced solution $u_N(\mu)$ and $s_N(\mu)$ are computed rapidly. The operation count of the online phase is ideally completely independent of $\mathcal{N}$, and typically scales polynomially in N .

In view of this desired computational splitting, we see that step 1 above clearly belongs to the offline phase, while step 4 is part of the online phase. But step 2 and 3 can not clearly be assigned to either of the two phases as they require expensive but as well parameter-dependent operations. This is where the parameter-separability Def. 2.5 comes into play by suitably dividing steps 2 and 3. The crucial insight is that due to linearity of the problem, parameter-separability of a, f, l transfers to parameter-separability of $\mathbf{A}_N, \mathbf{f}_N, \mathbf{l}_N$.

Corollary 2.29 (Offline/Online Decomposition of $(P_N(\mu))$).

(Offline Phase:) After computation of a reduced basis $\Phi_N = \{\varphi_1, \dots, \varphi_N\}$ construct the parameter-independent component matrices and vectors

$$\begin{aligned} \mathbf{A}_{N,q} &:= (a_q(\varphi_j, \varphi_i))_{i,j=1}^N \in \mathbb{R}^{N \times N}, q = 1, \dots, Q_A, \\ \mathbf{f}_{N,q} &:= (f_q(\varphi_i))_{i=1}^N \in \mathbb{R}^N, q = 1, \dots, Q_f, \\ \mathbf{l}_{N,q} &:= (l_q(\varphi_i))_{i=1}^N \in \mathbb{R}^N, q = 1, \dots, Q_l. \end{aligned}$$

(Online Phase:) For a given $\mu \in \mathcal{P}$ evaluate the coefficient functions $\theta_q^a(\mu), \theta_q^f(\mu),$

$\theta_q^l(\mu)$ for q in suitable ranges and assemble the matrix and vectors

$$\mathbf{A}_N(\mu) = \sum_{q=1}^{Q_a} \theta_q^a(\mu) \mathbf{A}_{N,q}, \quad \mathbf{f}_N(\mu) = \sum_{q=1}^{Q_f} \theta_q^f(\mu) \mathbf{f}_{N,q}, \quad \mathbf{l}_N(\mu) = \sum_{q=1}^{Q_l} \theta_q^l(\mu) \mathbf{l}_{N,q},$$

which exactly results in the discrete system of Prop. 2.11, which then can be solved for $u_N(\mu)$ and $s_N(\mu)$.

Note, that the computation of $\mathbf{A}_{N,q}$ can be realized in a very simple way: Let the reduced basis vectors φ_j be expanded in the basis $\{\psi_i\}_{i=1}^{\mathcal{N}}$ of the discrete full problem by $\varphi_j = \sum_{i=1}^{\mathcal{N}} \varphi_{i,j} \psi_i$ with coefficient matrix

$$\Phi_N := (\varphi_{i,j})_{i,j=1}^{\mathcal{N},N} \in \mathbb{R}^{\mathcal{N} \times N}. \quad (25)$$

If the component matrices $\mathbf{A}_q := (a_q(\psi_j, \psi_i))_{i,j=1}^{\mathcal{N}}$ and the component vectors $\mathbf{l}_q := (l_q(\psi_i))_{i=1}^{\mathcal{N}}$, $\mathbf{f}_q := (f_q(\psi_i))_{i=1}^{\mathcal{N}}$ from the full problem are available, the computations for the reduced matrices and vectors reduce to matrix-vector operations:

$$\mathbf{A}_{N,q} = \Phi_N^T \mathbf{A}_q \Phi_N, \quad \mathbf{f}_{N,q} = \Phi_N^T \mathbf{f}_q, \quad \mathbf{l}_{N,q} = \Phi_N^T \mathbf{l}_q.$$

Concerning complexities, we realize that the offline phase scales in the order of $\mathcal{O}(N \cdot \mathcal{N}^2 + N \cdot \mathcal{N}(Q_f + Q_l) + N^2 \cdot \mathcal{N} Q_a)$, the dominating part being the snapshot computation. We see from Cor. 2.29 that the online phase is computable in $\mathcal{O}(N^2 \cdot Q_a + N \cdot (Q_f + Q_l) + N^3)$, in particular completely independent of $\mathcal{N}$. This is the reason why we notationally do not discriminate between the analytical and finite element solution: $\mathcal{N}$ can be chosen arbitrarily high, i.e. the discrete full solution can be chosen arbitrarily accurate, without affecting the computational complexity of the online phase. Certainly, in practice, a certain finite $\mathcal{N}$ must be chosen and then the reduced dimension N must be adapted. Here it is important to note, that the RB-approximation and error estimation procedure only is informative as long as N is not too large and the reduction error is dominating the finite element error. A too large choice of N for a fixed given $\mathcal{N}$ does not make sense. In the limit case $N = \mathcal{N}$ we would have RB-errors and estimators being zero, hence exactly reproduce the discrete solution, but not the analytical (Sobolev space) solution.

The computational separation can also be illustrated by a runtime diagram, cf. Fig. 5. Let $t_{\text{full}}, t_{\text{offline}}, t_{\text{online}}$ denote the computational time for the single computation of a solution of $(P(\mu))$, the offline and the single computation of the online phase of $(P_N(\mu))$. Assuming that these times are constant for all parameters, we obtain linear/affine relations of the overall computation time on the number of simulation requests k : The overall time for k full solutions is $t(k) := k \cdot t_{\text{full}}$, while the reduced model (including offline phase) requires $t_N(k) := t_{\text{offline}} + k \cdot t_{\text{online}}$. As noted earlier, the reduced model pays off as soon as sufficiently many, i.e. $k > k^* := \frac{t_{\text{offline}}}{t_{\text{full}} - t_{\text{online}}}$ simulation requests are expected.

As noted earlier, the *certification* by a-posteriori error bounds is an important topic. Hence, we will next address the offline/online decomposition of the a-posteriori

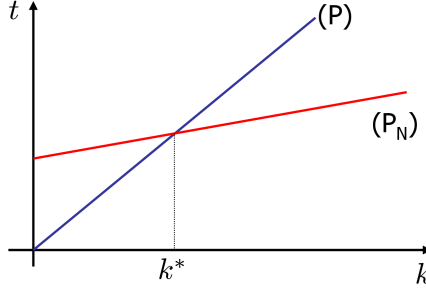


Figure 5: Runtime behavior of the full and the reduced model with increasing number k of simulations.

error estimators. The crucial insight is that parameter-separability also holds for the residuals and hence for the residual norms.

Proposition 2.30 (Parameter-Separability of the Residual). *Set $Q_r := Q_f + NQ_a$ and define $r_q \in X'$, $q = 1, \dots, Q_r$ via*

$$(r_1, \dots, r_{Q_r}) := (f_1, \dots, f_{Q_f}, a_1(\varphi_1, \cdot), \dots, a_{Q_a}(\varphi_1, \cdot), \dots, a_1(\varphi_N, \cdot), \dots, a_{Q_a}(\varphi_N, \cdot)).$$

Let $u_N(\mu) = \sum_{i=1}^N u_{N,i} \varphi_i$ be the solution of $(P_N(\mu))$ and define $\theta_q^r(\mu)$, $q = 1, \dots, Q_r$ by

$$(\theta_1^r, \dots, \theta_{Q_r}^r) := (\theta_1^f, \dots, \theta_{Q_f}^f, -\theta_1^a \cdot u_{N,1}, \dots, -\theta_{Q_a}^a \cdot u_{N,1}, \dots, -\theta_1^a \cdot u_{N,N}, \dots, -\theta_{Q_a}^a \cdot u_{N,N}).$$

Let $v_r, v_{r,q} \in X$ denote the Riesz-representatives of r, r_q , i.e. $r(v) = \langle v_r, v \rangle$ and $r_q(v) = \langle v_{r,q}, v \rangle$, $v \in X$. Then the residual and its Riesz-representatives are parameter-separable via

$$r(v; \mu) = \sum_{q=1}^{Q_r} \theta_q^r(\mu) r_q(v), \quad v_r(\mu) = \sum_{q=1}^{Q_r} \theta_q^r(\mu) v_{r,q}, \quad \mu \in \mathcal{P}, v \in X. \quad (26)$$

Proof. Linearity of $a(\cdot, \cdot; \mu)$ directly implies that the first equation in (26) is a reformulation of r defined in (6). Linearity of the Riesz-map gives the second statement in (26) for the Riesz-representative. $\square$

In the error estimation procedure, it is necessary to compute Riesz-representatives of linear functionals. We briefly want to comment on how this can be realized.

Lemma 2.31 (Computation of Riesz-representatives). *Let $g \in X'$ and $X = \text{span}(\phi_i)_{i=1}^{\mathcal{N}}$ with basis functions ϕ_i . We introduce the coefficient vector $\mathbf{v} = (v_i)_{i=1}^{\mathcal{N}} \in \mathbb{R}^{\mathcal{N}}$ of the*

Riesz-representative $v_g = \sum_{i=1}^{\mathcal{N}} v_i \psi_i \in X$. Then, $\mathbf{v}$ can simply be obtained by solving the linear system

$$\mathbf{K}\mathbf{v} = \mathbf{g}$$

with vector $\mathbf{g} = (g(\psi_i))_{i=1}^{\mathcal{N}} \in \mathbb{R}^{\mathcal{N}}$ and (typically sparse) inner product matrix $\mathbf{K}$ given in (22).

Proof. We verify for any test function $u = \sum_{i=1}^{\mathcal{N}} u_i \psi_i$ with coefficient vector $\mathbf{u} = (u_i)_{i=1}^{\mathcal{N}} \in \mathbb{R}^{\mathcal{N}}$

$$g(u) = \sum_{i=1}^{\mathcal{N}} u_i g(\psi_i) = \mathbf{u}^T \mathbf{g} = \mathbf{u}^T \mathbf{K} \mathbf{v} = \left\langle \sum_{i=1}^{\mathcal{N}} u_i \psi_i, \sum_{j=1}^{\mathcal{N}} v_j \psi_j \right\rangle = \langle v_g, u \rangle.$$

□

The parameter-separability of the residual allows to compute the norm of the residual in an offline/online decomposition.

Proposition 2.32 (Offline/Online Decomposition of the Residual Norm).

(Offline Phase:) After the offline phase of the RB-model according to Cor. 2.29 we define the matrix

$$\mathbf{G}_r := (r_q(v_{r,q'}))_{q,q'=1}^{Q_r} \in \mathbb{R}^{Q_r \times Q_r}$$

via the residual components r_q and their Riesz-representatives $v_{r,q}$.

(Online Phase:) For given μ and RB-solution $u_N(\mu)$, we compute the residual coefficient vector $\theta_r(\mu) := (\theta_1^r(\mu), \dots, \theta_{Q_r}^r(\mu))^T \in \mathbb{R}^{Q_r}$ and obtain

$$\|r(\cdot; \mu)\|_{X'} = \sqrt{\theta_r(\mu)^T \mathbf{G}_r \theta_r(\mu)}.$$

Proof. First we realize that $\mathbf{G}_r = (\langle v_{r,q}, v_{r,q'} \rangle)_{q,q'=1}^{Q_r}$ due to definition of the Riesz-representatives. Isometry of the Riesz-map and parameter-separability (26) yield

$$\|r(\mu)\|_{X'}^2 = \|v_r(\mu)\|^2 = \left\langle \sum_{q=1}^{Q_r} \theta_q^r(\mu) v_{r,q}, \sum_{q'=1}^{Q_r} \theta_{q'}^r(\mu) v_{r,q'} \right\rangle = \theta_r(\mu)^T \mathbf{G}_r \theta_r(\mu).$$

□

Again, the online phase is independent of $\mathcal{N}$ as it has complexity $\mathcal{O}(Q_r^2)$. Completely analogous, we can compute the dual norm of the output functional used in the output bound (9); we omit the proof:

Proposition 2.33 (Offline/Online Decomposition of $\|l(\cdot; \mu)\|_{X'}$).

(Offline Phase:) We compute the Riesz-representatives $v_{l,q} \in X$ of the output functional components, i.e. $\langle v_{l,q}, v \rangle = l_q(v)$, $v \in X$ and define the matrix

$$\mathbf{G}_l := (l_q(v_{l,q'}))_{q,q'=1}^{Q_l} \in \mathbb{R}^{Q_l \times Q_l}.$$

(Online Phase:) For given μ compute the functional coefficient vector

$$\theta_l(\mu) := (\theta_1^l(\mu), \dots, \theta_{Q_l}^l(\mu))^T \in \mathbb{R}^{Q_l}$$

and obtain

$$\|l(\cdot; \mu)\|_{X'} = \sqrt{\theta_l(\mu)^T \mathbf{G}_l \theta_l(\mu)}.$$

A further quantity appearing in relative a-posteriori error estimates is the norm of u_N , which can be similarly decomposed.

Proposition 2.34 (Offline/Online Decomposition of $\|u_N(\mu)\|$).

(Offline Phase:) After the offline phase of the RB-model according to Cor. 2.29 we define the reduced inner product matrix

$$\mathbf{K}_N := (\langle \varphi_i, \varphi_j \rangle)_{i,j=1}^N \in \mathbb{R}^{N \times N}. \quad (27)$$

(Online Phase:) For given μ and $u_N(\mu) \in \mathbb{R}^N$ computed in the online phase according to Cor. 2.29 we compute

$$\|u_N(\mu)\| = \sqrt{\mathbf{u}_N^T(\mu) \mathbf{K}_N \mathbf{u}_N(\mu)}.$$

Proof. We directly verify that

$$\|u_N\|^2 = \left\langle \sum_i u_{N,i} \varphi_i, \sum_j u_{N,j} \varphi_j \right\rangle = \sum_{i,j=1}^N u_{N,i} u_{N,j} \langle \varphi_i, \varphi_j \rangle = \mathbf{u}_N^T \mathbf{K}_N \mathbf{u}_N.$$

□

Here, too, the online phase is independent of $\mathcal{N}$ as it has complexity $\mathcal{O}(N^2)$. Again, $\mathbf{K}_N$ can be easily obtained by matrix operations. With the full inner product matrix $\mathbf{K}$ defined in (22) and the RB-coefficient matrix Φ_N from (25) we compute

$$\mathbf{K}_N = \Phi_N^T \mathbf{K} \Phi_N.$$

Analogously, for the relative energy norm bound in the symmetric case, the energy norm can be computed as

$$\|u_N\|_\mu^2 = \mathbf{u}_N^T \mathbf{A}_N \mathbf{u}_N.$$

The remaining ingredient of the a-posteriori error estimators is the computation of lower bounds $\alpha_{\text{LB}}(\mu)$ of the coercivity constant. We note that using the uniform lower bound is a viable choice, i.e. $\alpha_{\text{LB}}(\mu) := \tilde{\alpha}$, if the latter is available/computable a-priori. Further, for some model problems, $\alpha(\mu)$ may be exactly and rapidly computable, hence $\alpha_{\text{LB}}(\mu) := \alpha(\mu)$ may be a valid choice. For example, this is indeed available for the thermal block model, cf. Exercise 5.11.

A more general approach, the so called *min-theta* procedure [57], can be applied under certain assumptions. It makes use of the parameter-separability and the (expensive, offline) computation of a single coercivity constant for the full problem. Then this lower bound can be evaluated rapidly in the online phase.

Proposition 2.35 (Min-Theta Approach for Computing $\alpha_{\text{LB}}(\mu)$). *We assume that the components of $a(\cdot, \cdot; \mu)$ satisfy $a_q(u, u) \geq 0, q = 1, \dots, Q_a, u \in X$ and the coefficient functions fulfill $\theta_q^a(\mu) > 0, \mu \in \mathcal{P}$. Let $\bar{\mu} \in \mathcal{P}$ such that $\alpha(\bar{\mu})$ is available. Then we have*

$$0 < \alpha_{\text{LB}}(\mu) \leq \alpha(\mu), \quad \mu \in \mathcal{P}$$

with the lower bound

$$\alpha_{\text{LB}}(\mu) := \alpha(\bar{\mu}) \cdot \min_{q=1, \dots, Q_a} \frac{\theta_q^a(\mu)}{\theta_q^a(\bar{\mu})}.$$

Proof. As $0 < \alpha(\bar{\mu})$ and $0 < C(\mu) := \min_q \theta_q^a(\mu)/\theta_q^a(\bar{\mu})$, we have $0 < \alpha(\bar{\mu})C(\mu) = \alpha_{\text{LB}}(\mu)$. For all $u \in X$ holds

$$\begin{aligned} a(u, u; \mu) &= \sum_{q=1}^{Q_a} \theta_q^a(\mu) a_q(u, u) = \sum_{q=1}^{Q_a} \frac{\theta_q^a(\mu)}{\theta_q^a(\bar{\mu})} \theta_q^a(\bar{\mu}) a_q(u, u) \\ &\geq \sum_{q=1}^{Q_a} \left(\min_{q'=1, \dots, Q_a} \frac{\theta_q^a(\mu)}{\theta_{q'}^a(\bar{\mu})} \right) \theta_q^a(\bar{\mu}) a_q(u, u) \\ &= C(\mu) a(u, u; \bar{\mu}) \geq C(\mu) \alpha(\bar{\mu}) \|u\|^2 = \alpha_{\text{LB}}(\mu) \|u\|^2. \end{aligned}$$

In particular $\alpha(\mu) = \inf_{u \in X \setminus \{0\}} (a(u, u; \mu) / \|u\|^2) \geq \alpha_{\text{LB}}(\mu)$. $\square$

This lower bound is obviously computable in $\mathcal{O}(Q_a)$, hence fast, if we assume a small/decent number of components Q_a .

For the *min-theta* approach, we require a single evaluation of $\alpha(\mu)$ for $\mu = \bar{\mu}$ in the offline phase. This can be obtained via solving a high-dimensional, hence expensive, eigenvalue problem, cf. [57] for the continuous formulation.

Proposition 2.36 (Computation of $\alpha(\mu)$ of Discretized Full Problem). *Let $\mathbf{A}(\mu), \mathbf{K} \in \mathbb{R}^{\mathcal{N} \times \mathcal{N}}$ denote the high-dimensional discrete system and inner product matrix as given in (22). Define $\mathbf{A}_s(\mu) := \frac{1}{2}(\mathbf{A}(\mu) + \mathbf{A}^T(\mu))$ as the symmetric part of $\mathbf{A}(\mu)$. Then*

$$\alpha(\mu) = \lambda_{\min}(\mathbf{K}^{-1} \mathbf{A}_s(\mu)), \quad (28)$$

where $\lambda_{\min}$ denotes the smallest eigenvalue.

Proof. We make use of a decomposition $\mathbf{K} = \mathbf{L}\mathbf{L}^T$, e.g. Cholesky or matrix square root, a substitution $\mathbf{v} := \mathbf{L}^T \mathbf{u}$, and omit the parameter for notational simplicity:

$$\alpha = \inf_{u \in X} \frac{a(u, u)}{\|u\|^2} = \inf_{u \in \mathbb{R}^{\mathcal{N}}} \frac{u^T \mathbf{A} u}{u^T \mathbf{K} u} = \inf_{u \in \mathbb{R}^{\mathcal{N}}} \frac{u^T \mathbf{A}_s u}{u^T \mathbf{K} u} = \inf_{v \in \mathbb{R}^{\mathcal{N}}} \frac{v^T \mathbf{L}^{-1} \mathbf{A}_s \mathbf{L}^{-T} v}{v^T v}.$$

Thus α is a minimum of a Rayleigh-quotient, hence the smallest eigenvalue of the symmetric matrix $\tilde{\mathbf{A}}_s := \mathbf{L}^{-1} \mathbf{A}_s \mathbf{L}^{-T}$. The matrices $\tilde{\mathbf{A}}_s$ and $\mathbf{K}^{-1} \mathbf{A}_s$ are similar as

$$\mathbf{L}^T (\mathbf{K}^{-1} \mathbf{A}_s) \mathbf{L}^{-T} = \mathbf{L}^T \mathbf{L}^{-T} \mathbf{L}^{-1} \mathbf{A}_s \mathbf{L}^{-T} = \tilde{\mathbf{A}}_s,$$

hence they have identical eigenvalues, which proves (28). $\square$

Certainly, an inversion of $\mathbf{K}$ needs to be prevented, hence in practice one can use an eigenvalue solver which only requires matrix-vector products: As soon as a product $\mathbf{y} = \mathbf{K}^{-1}\mathbf{A}_s\mathbf{x}$ is required, one solves the system $\mathbf{K}\mathbf{y} = \mathbf{A}_s\mathbf{x}$. Alternatively, one can make use of solvers for generalized eigenvalue problems of the form $\mathbf{A}_s\mathbf{u} = \lambda\mathbf{K}\mathbf{u}$ and determine the smallest generalized eigenvalue.

Concerning computational complexity for the a-posteriori error estimators $\Delta_u(\mu)$ and $\Delta_s(\mu)$ we obtain as offline-complexity $\mathcal{O}(\mathcal{N}^3 + \mathcal{N}^2(Q_f + Q_l + NQ_a) + \mathcal{N}Q_l^2)$. The dominating part corresponds to an eigenvalue problem in cubic complexity for the computation of $\alpha(\bar{\mu})$ for the min-theta-procedure. Then, the online phase merely scales as $\mathcal{O}((Q_f + NQ_a)^2 + Q_l^2 + Q_a)$, again fully independent of the high dimension $\mathcal{N}$.

There also exist techniques for computing upper bounds $\gamma_{\text{UB}}(\mu)$ of continuity constants $\gamma(\mu)$, cf. Exercise 5.12 for a *max-theta* approach. This allows to evaluate effectivity bounds according to (10), etc. online.

For problems where the min-theta approach cannot be applied, the so called Successive Constraint Method (SCM) [40] is an option. Based on precomputation of many $\alpha(\mu^{(i)})$ in the offline phase, in the online phase a small linear optimization problem is solved for any new $\mu \in \mathcal{P}$, which gives a rigorous lower bound $\alpha_{\text{LB}}(\mu)$.

Definition 2.37 (Successive Constraint Method (SCM)). *Let $a(\cdot, \cdot; \mu)$ be uniformly coercive with respect to μ and parameter separable with $Q := Q_a$ components. Let $C, D \subset \mathcal{P}$ be finite subsets and $M_-, M_+ \in \mathbb{N}$. Define*

$$Y := \{y = (y_1, \dots, y_Q) \in \mathbb{R}^Q \mid \exists u \in X \text{ with } y_q = a_q(u, u)/\|u\|^2, q = 1, \dots, Q\}.$$

We define a target function $J : \mathcal{P} \times \mathbb{R}^Q \rightarrow \mathbb{R}$ by

$$J(\mu, y) := \sum_{q=1}^Q \theta_q^a(\mu) y_q$$

and a polytope B_Q by

$$\sigma_q^- := \inf_{u \in X} \frac{a_q(u, u)}{\|u\|^2}, \quad \sigma_q^+ := \sup_{u \in X} \frac{a_q(u, u)}{\|u\|^2}, \quad (29)$$

$$B_Q := \prod_{q=1}^Q [\sigma_q^-, \sigma_q^+] \subset \mathbb{R}^Q. \quad (30)$$

For $M \in \mathbb{N}, \mu \in \mathcal{P}$ define $P_M(\mu, C) \subset C$ by

$$P_M(\mu, C) := \begin{cases} M - \text{nearest neighbours of } \mu \text{ in } C & \text{if } 1 \leq M \leq |C| \\ C & \text{if } |C| \leq M \\ \emptyset & \text{if } M = 0. \end{cases}$$

Then, we define for $\mu \in \mathcal{P}$ the sets

$$\begin{aligned} Y_{\text{LB}}(\mu) &:= \{y \in B_Q \mid J(\mu', y) \geq \alpha(\mu') \forall \mu' \in P_{M_x}(\mu, C) \text{ and} \\ &\quad J(\mu', y) \geq 0 \forall \mu' \in P_{M_+}(\mu, D)\} \\ Y_{\text{UB}} &:= \{y^*(\mu') \mid \mu' \in C\} \quad \text{with} \quad y^*(\mu') := \arg \min_{y \in Y} J(\mu', y) \end{aligned}$$

and herewith the quantities

$$\alpha_{\text{LB}}(\mu) := \min_{y \in Y_{\text{LB}}(\mu)} J(\mu, y), \quad \alpha_{\text{UB}}(\mu) := \min_{y \in Y_{\text{UB}}} J(\mu, y). \quad (31)$$

With the above definitions, it is easy to show the bounding property:

Proposition 2.38 (Coercivity Constant Bounds by SCM). *For all $\mu \in \mathcal{P}$ holds*

$$\alpha_{\text{LB}}(\mu) \leq \alpha(\mu) \leq \alpha_{\text{UB}}(\mu). \quad (32)$$

Proof. First we see that

$$\alpha(\mu) = \inf_{u \in X \setminus \{0\}} \frac{\sum_{q=1}^Q \theta_q^a(\mu) a_q(u, u)}{\|u\|^2} = \min_{y \in Y} J(\mu, y). \quad (33)$$

Further we see that $Y_{\text{UB}} \subset Y \subset Y_{\text{LB}}(\mu)$. For the first inclusion we verify for any $y \in Y_{\text{UB}}$ that there exists $\mu' \in C$ with

$$y = y^*(\mu') = \arg \min_{\tilde{y} \in Y} J(\mu', \tilde{y}),$$

thus $y \in Y$. For the second inclusion we choose $y \in Y$, thus $y \in B_Q$ and see for any $\mu' \in C$

$$\alpha(\mu') = \min_{\tilde{y} \in Y} J(\mu', \tilde{y}) \leq J(\mu', y).$$

Analogously, for any $\mu' \in D$

$$0 < \alpha(\mu') \leq J(\mu', y).$$

Hence, $y \in Y_{\text{LB}}(\mu)$. The nestedness of the sets then yields

$$\min_{y \in Y_{\text{LB}}(\mu)} J(\mu, y) \leq \min_{y \in Y} J(\mu, y) \leq \min_{y \in Y_{\text{UB}}} J(\mu, y)$$

which directly implies (32) with the definition of the bounds (31) and (33). $\square$

For a fixed μ the function J is obviously indeed linear in y , which implies the necessity to solve of a small linear optimization problem in the online phase. For further details on the SCM, we refer to [40, 63].

This concludes the section as we provided offline/online computational procedures to evaluate all ingredients required for efficient computation of the reduced solution, a-posteriori error estimates and effectivity bounds.

2.6 Basis Generation

In this section, we address the topic of basis generation. While this point can be seen as a part of the offline phase for generation of the reduced model, it is such a central issue that we devote this separate section to it. One may ask why the basis generation was not presented prior to the RB-methodology of the previous sections. The reason is that the basis generation will make full use of the presented tools of Sec. 2.3 and 2.5 for constructively generating a problem-dependent reduced basis.

We already used the most simple reduced basis type earlier:

Definition 2.39 (Lagrangian Reduced Basis). *Let $S_N := \{\mu^{(1)}, \dots, \mu^{(N)}\} \subset \mathcal{P}$ be such that the snapshots $\{u(\mu^{(i)})\}_{i=1}^N \subset X$ are linearly independent. We then call*

$$\Phi_N := \{u(\mu^{(1)}), \dots, u(\mu^{(N)})\}$$

a Lagrangian reduced basis.

An alternative to a Lagrangian Reduced basis may be seen in a Taylor reduced basis [21] where one includes sensitivity derivatives of the solution around a certain parameter, e.g. a first order *Taylor reduced basis*

$$\Phi_N := \{u(\mu^{(0)}), \partial_{\mu_1} u(\mu^{(0)}), \dots, \partial_{\mu_p} u(\mu^{(0)})\}.$$

While this basis is expected to give rather good approximation locally around $\mu^{(0)}$, a Lagrangian reduced basis has the ability to provide globally well-approximating models if S_N is suitably chosen, cf. the interpolation argument of Rem. 2.18.

As we are aiming at global approximation of the manifold $\mathcal{M}$ we define a corresponding error measure. We would be interested in finding a space X_N of dimension N minimizing

$$E_N := \sup_{\mu \in \mathcal{P}} \|u(\mu) - u_N(\mu)\|. \quad (34)$$

This optimization over all subspaces of a given dimension N is a very complex optimization problem. Hence, a practical relaxation is an incremental procedure: Construct an approximating subspace by iteratively adding new basis vectors. The choice of each new basis vector is led by the aim of minimizing E_N . This is the rationale behind the *greedy procedure*, which was first used in an RB-context in [71] and meanwhile is standard for stationary problems. It incrementally constructs both the sample set S_N and the basis Φ_N . We formulate the abstract algorithm and comment on practical aspects and choices for its realization. As main ingredient we require an error indicator $\Delta(Y, \mu) \in \mathbb{R}^+$ that predicts the expected approximation error for the parameter μ when using $X_N = Y$ as approximation space.

Definition 2.40 (Greedy Procedure). *Let $S_{\text{train}} \subset \mathcal{P}$ be a given training set of parameters and $\varepsilon_{\text{tol}} > 0$ a given error tolerance. Set $X_0 := \{0\}$, $S_0 = \emptyset$, $\Phi_0 := \emptyset$, $n := 0$ and define*

iteratively

$$\begin{aligned}
\text{while } \varepsilon_n &:= \max_{\mu \in S_{\text{train}}} \Delta(X_n, \mu) > \varepsilon_{\text{tol}} & (35) \\
\mu^{(n+1)} &:= \arg \max_{\mu \in S_{\text{train}}} \Delta(X_n, \mu) \\
S_{n+1} &:= S_n \cup \{\mu^{(n+1)}\} \\
\varphi_{n+1} &:= u(\mu^{(n+1)}) \\
\Phi_{n+1} &:= \Phi_n \cup \{\varphi_{n+1}\} \\
X_{n+1} &:= X_n \oplus \text{span}(\varphi_{n+1}) \\
n &\leftarrow n + 1 \\
&\text{end while.}
\end{aligned}$$

The algorithm produces the desired RB-space X_N and basis Φ_N by setting $N := n + 1$ as soon as (35) is false. We can state a simple termination criterion for the above algorithm: If for all $\mu \in \mathcal{P}$ and subspaces $Y \subset X$ holds that

$$u(\mu) \in Y \Rightarrow \Delta(Y, \mu) = 0, \quad (36)$$

then the above algorithm terminates in at most $N \leq |S_{\text{train}}|$ steps, where $|\cdot|$ indicates the cardinality of a given set. The reason is that with (36) no sample in S_{train} will be selected twice. This criterion will be easily satisfied by reasonable indicators.

Alternatively, the first iteration, i.e. determination of $\mu^{(1)}$ is frequently skipped by choosing a random initial parameter vector.

The greedy procedure generates a Lagrangian reduced basis with a carefully selected sample set. The basis is *hierarchical* in the sense that $\Phi_n \subset \Phi_m$ for $n \leq m$. This allows to adjust the accuracy of the reduced model online by varying its dimension.

The training set S_{train} is mostly chosen as a (random or structured) finite subset. The maximization then is a linear search. The training set must represent $\mathcal{P}$ well in order to not “miss” relevant parts of the parameter domain. In practice it should be taken as large as possible.

Remark 2.41 (Choice of Error Indicator $\Delta(Y, \mu)$). *There are different options for the choice of the error indicator $\Delta(Y, \mu)$ in the greedy procedure, each with advantages, disadvantages and requirements.*

- i) *Projection error as indicator: In some cases, it can be recommended to use the above algorithm with the best-approximation error (which is an orthogonal projection error)*

$$\Delta(Y, \mu) := \inf_{v \in Y} \|u(\mu) - v\| = \|u(\mu) - P_Y u(\mu)\|.$$

Here, P_Y denotes the orthogonal projection onto Y . In this version of the greedy procedure, the error indicator is expensive to evaluate as high-dimensional operations are required, hence S_{train} must be of moderate size. Also, all snapshots $u(\mu)$ must be available, which possibly limits the size of S_{train} due to memory constraints.

Still the advantage of this approach is that the RB-model is decoupled from the basis generation: No RB-model or a-posteriori error estimators are required, the algorithm purely constructs a good approximation space. By statements such as the best-approximation relation, Prop. 2.14, one can then be sure that the corresponding RB-model using the constructed X_N will be good. This version of the greedy algorithm will be denoted strong greedy procedure.

- ii) True RB-error as indicator: If one has an RB-model available, but no a-posteriori error bounds, one can use

$$\Delta(Y, \mu) := \|u(\mu) - u_N(\mu)\|.$$

Again, in this version of the greedy procedure, the error indicator is expensive to evaluate, hence S_{train} must be of moderate size. And again, all snapshots $u(\mu)$ must be available, limiting the size of S_{train} . The advantage of this approach is that the error criterion which is minimized exactly is the measure used in E_N in (34).

- iii) A-posteriori error estimator as indicator: This is the recommended approach, if one has both an RB-model and an a-posteriori error estimator available. Hence we choose

$$\Delta(Y, \mu) := \Delta_u(\mu).$$

The evaluation of $\Delta(Y, \mu) = \Delta_u(\mu)$ (or a relative estimator) is very cheap, hence the sample set S_{train} can be chosen much larger than when using a true RB- or projection-error as indicator. By this, the training set S_{train} can be expected to be much more representative for the complete parameter space in contrast to a smaller training set. No snapshots need to be precomputed. In the complete greedy procedure only N high-dimensional solves of $(P(\mu))$ are required. Hence, the complete greedy procedure is expected to be rather fast. This version of the greedy algorithm is called a weak greedy, as will be explained more precisely in the subsequent convergence analysis.

Note that all of these choices for $\Delta(Y, \mu)$ satisfy (36): This statement is trivial for i) the projection error. For the RB-error ii) it is a consequence of the reproduction of solutions, Prop. 2.16, and for the a-posteriori error estimators, it is a consequence of the vanishing error bound, Cor. 2.20. Hence the greedy algorithm is guaranteed to terminate.

Alternatively, one can also use goal-oriented indicators, i.e. $\Delta(Y, \mu) = |s(\mu) - s_N(\mu)|$ or $\Delta_s(\mu)$. One can expect to obtain a rather small basis which approximates $s(\mu)$ very well, but $u(\mu)$ will possibly not be well approximated. In contrast, by using the above indicators i), ii), iii), one can expect to obtain a larger basis, which accurately approximates $u(\mu)$ as well as the output $s(\mu)$.

Note, that in general one cannot expect monotonical decay of ε_n for $n = 1, \dots, N$. Only in certain cases this can be proven, cf. Exercise 5.13.

Remark 2.42 (Overfitting, quality measurement). The error sequence $\{\varepsilon_n\}_{n=0}^N$ generated by the greedy procedure in (35) is only a training error in statistical learning terminology. The quality of a model, its generalization capabilities, can not necessarily be

concluded from this due to possible overfitting. This means that possibly

$$\max_{\mu \in \mathcal{P}} \Delta(X_N, \mu) \gg \varepsilon_N.$$

If ε_{tol} is sufficiently small, for example, we obtain $N = |S_{\text{train}}|$ and we will have the training error $\varepsilon_N = 0$. But the model is very likely not exact on the complete parameter set. Hence, it is always recommended to evaluate the quality of a model on an independent test set $S_{\text{test}} \subset \mathcal{P}$, which is not related to S_{train} .

We give some hints on theoretical foundation of the above greedy procedure. Until recently, this algorithm seemed to be a heuristic procedure that works very well in practice in various cases. Rigorous analysis was not available. But then a useful approximation-theoretic result has been formulated, first concerning exponential convergence [9], then also for algebraic convergence [5]. It states that if $\mathcal{M}$ can be approximated well by some linear subspace, then the greedy algorithm will identify approximation spaces which are only slightly worse than these optimal subspaces. The optimal subspaces are defined via the *Kolmogorov n -width* defined as the maximum error of the best-approximating linear subspace

$$d_n(\mathcal{M}) := \inf_{\substack{Y \subset X \\ \dim Y = n}} \sup_{u \in \mathcal{M}} \|u - P_Y u\|. \quad (37)$$

The convergence statement [5] adopted to our notation and assumptions then can be formulated as follows; we omit the proof.

Proposition 2.43 (Greedy Convergence Rates). *Let $S_{\text{train}} = \mathcal{P}$ be compact, and the error indicator Δ chosen such that for suitable $\gamma \in (0, 1]$ holds*

$$\|u(\mu^{(n+1)}) - P_{X_n} u(\mu^{(n+1)})\| \geq \gamma \sup_{u \in \mathcal{M}} \|u - P_{X_n} u\|. \quad (38)$$

i) (Algebraic convergence rate:) if $d_n(\mathcal{M}) \leq M n^{-\alpha}$ for some $\alpha, M > 0$ and all $n \in \mathbb{N}$ and $d_0(\mathcal{M}) \leq M$ then

$$\varepsilon_n \leq C M n^{-\alpha}, n > 0$$

with a suitable (explicitly computable) constant $C > 0$.

ii) (Exponential convergence rate:) if $d_n(\mathcal{M}) \leq M e^{-an^a}$ for $n \geq 0, M, a, \alpha > 0$ then

$$\varepsilon_n \leq C M e^{-cn^\beta}, n \geq 0$$

with $\beta := \alpha/(\alpha + 1)$ and suitable (explicitly computable) constants $c, C > 0$.

Remark 2.44 (Strong versus Weak Greedy). If $\gamma = 1$ (e.g. obtained for the choice 2.41i), $\Delta(Y, \mu) := \|u(\mu) - P_Y u(\mu)\|$ the algorithm is called a strong greedy, while for $\gamma < 1$ the algorithm is called weak greedy algorithm. Note, that (38) is valid in the case of $\Delta(Y, \mu) := \Delta_\mu(\mu)$, i.e. 2.41iii), due to the Lemma of Céa (Prop. 2.14), the effectivity (Prop. 2.21) and the error bound property (Prop. 2.19). Using the notation $u_N(\mu), \Delta_\mu(\mu)$

for the RB-solution and estimator using the corresponding intermediate spaces X_n for $1 \leq n \leq N$ we derive

$$\begin{aligned}
& \left\| u(\mu^{(n+1)}) - P_{X_n} u(\mu^{(n+1)}) \right\| = \inf_{v \in X_n} \left\| u(\mu^{(n+1)}) - v \right\| \\
& \geq \frac{\alpha(\mu)}{\gamma(\mu)} \left\| u(\mu^{(n+1)}) - u_N(\mu^{(n+1)}) \right\| \geq \frac{\alpha(\mu)}{\gamma(\mu)\eta_u(\mu)} \Delta_u(\mu^{(n+1)}) \\
& = \frac{\alpha(\mu)}{\gamma(\mu)\eta_u(\mu)} \sup_{\mu \in \mathcal{P}} \Delta_u(\mu) \geq \frac{\alpha(\mu)}{\gamma(\mu)\eta_u(\mu)} \sup_{\mu \in \mathcal{P}} \|u(\mu) - u_N(\mu)\| \\
& \geq \frac{\alpha(\mu)}{\gamma(\mu)\eta_u(\mu)} \sup_{\mu \in \mathcal{P}} \|u(\mu) - P_{X_n} u(\mu)\| \geq \frac{\tilde{\alpha}^2}{\tilde{\gamma}^2} \sup_{\mu \in \mathcal{P}} \|u(\mu) - P_{X_n} u(\mu)\|.
\end{aligned}$$

Hence, a weak greedy algorithm with parameter $\gamma := \frac{\tilde{\alpha}^2}{\tilde{\gamma}^2} \in (0, 1]$ is obtained.

As a result the greedy algorithm is theoretically well founded. Now merely the question arises, when “good approximability” is to be expected for a given problem. A positive answer has been given in [52, 57]. The assumptions in the statement are, for example, satisfied for the thermal block with $B_1 = 2, B_2 = 1$, fixing $\mu_1 = 1$ and choosing a single scalar parameter $\mu := \mu_2$.

Proposition 2.45 (Global Exponential Convergence for $p = 1$). *Let $\mathcal{P} = [\mu_{\min}, \mu_{\max}] \subset \mathbb{R}^+$ with $0 < \mu_{\min} < 1, \mu_{\max} = 1/\mu_{\min}$. Further, assume that $a(u, v; \mu) := \mu a_1(u, v) + a_2(u, v)$ is symmetric, f is not parameter-dependent, $a := \ln \frac{\mu_{\max}}{\mu_{\min}} > \frac{1}{2e}$ and $N_0 := 1 + \lfloor 2ea + 1 \rfloor$. For $N \in \mathbb{N}$ define S_N via $\mu_{\min} = \mu^{(1)} < \dots < \mu^{(N)} = \mu_{\max}$ with logarithmically equidistant samples and X_N the corresponding Lagrangian RB-space. Then*

$$\frac{\|u(\mu) - u_N(\mu)\|_{\mu}}{\|u(\mu)\|_{\mu}} \leq e^{-\frac{N-1}{N_0-1}}, \mu \in \mathcal{P}, N \geq N_0.$$

With uniform boundedness of the solution and norm-equivalence we directly obtain the same rate (just with an additional constant factor) for the error $\|u(\mu) - u_N(\mu)\|$.

We proceed with further aspects concerning computational procedures.

Remark 2.46 (Training Set Treatment). *There are several ways, how the training set can be treated slightly differently, leading to improvements.*

i) *Multistage Greedy: The first approach aims at a runtime acceleration: Instead of working with a fixed large training set S_{train} , which gives rise to $\mathcal{O}(|S_{\text{train}}|)$ runtime complexity in the greedy algorithm, one generates coarser subsets of this large training set:*

$$S_{\text{train}}^{(0)} \subset S_{\text{train}}^{(1)} \subset \dots \subset S_{\text{train}}^{(m)} := S_{\text{train}}.$$

Then, the greedy algorithm is started on $S_{\text{train}}^{(0)}$ resulting in a basis $\Phi_{N^{(0)}}$. This basis is used as starting basis for the greedy algorithm on the next larger training set. And this

procedure is repeated until the greedy is run on the complete training set, but with a large starting basis $\Phi_{N^{(m-1)}}$. The rationale behind this procedure is that many iterations will be performed on small training sets, while the few final iterations still guarantee the precision on the complete large training set. Overall, a remarkable runtime improvement can be obtained, while the quality of the basis is not expected to degenerate too much. Such an approach has been introduced as multistage greedy procedure [65].

ii) *Training set adaptation*: The next procedure aims at adaptation of the training set in order to realize uniform error distribution over the parameter space. For a given problem it is not clear a-priori, how the training set should be chosen best. If the training set is chosen too large, the offline runtime may be too high. If the training set is too small, overfitting may easily be obtained, cf. Rem. 2.42. The idea of the adaptive training set refinement [29, 30] is to start the greedy with a coarse set of training parameters, which are vertices of a mesh on the parameter domain. Then, in the finite element spirit, a-posteriori error estimators for subdomains are evaluated, grid cells with large error are marked for refinement, the marked cells are refined and the new vertices added to the training set. By this procedure, the training set is adapted to the problem at hand. The procedure may adaptively identify “difficult” parameter regions for example small diffusion constant values and refine more in such regions.

iii) *Randomization*: A simple idea allows to implicitly work with a large training set: When working with randomly drawn parameter samples, one can draw a new set S_{train} in each greedy loop iteration. Hereby, the effective parameter set that is involved in the training is virtually enlarged by a factor N . This idea and refinements have been presented in [37].

iv) *Full optimization*: In special cases, a true (local) optimization over the parameter space in (35), i.e. $S_{\text{train}} = \mathcal{P}$ can be realized, too [68]. The choice of a large training set is then reduced to a choice of a small set of multiple starting points for the highly nonlinear optimization procedure.

Remark 2.47 (Parameter Domain Partitioning). *The greedy procedure allows to prescribe accuracy via ε_{tol} and obtain a basis of size N that is a-priori unpredictable and hence the final online runtime is unclear. It would be desirable to control both the accuracy (by prescribing ε_{tol}) as well as the online runtime (by demanding $N \leq N_{\text{max}}$). The main idea to obtain this is via parameter domain partitioning:*

i) *hp-RB-approach* [20, 19]: Based on adaptive bisection of the parameter domain into subdomains, a partitioning of the parameter domain is generated (*h*-adaptivity). Then, small local bases can be generated for the different subdomains. If the accuracy and basis size criterion are not both satisfied for a subdomain, this subdomain is again refined and bases on the subdomains are generated. Finally, one has a collection of problems of type $(P_N(\mu))$, where $\mathcal{P}$ now is reduced to each of the subdomains of the partitioning. For a newly given μ in the online phase, merely the correct subdomain and model need to be identified by a search in the grid hierarchy. This method balances offline cost, both in terms of computational and storage requirements, against online accuracy and runtime.

ii) *P-partition*: A variant of parameter domain partitioning using hexaedral partitioning of the parameter space guarantees shape-regularity of the subdomains [29]. The method prevents partitioning into long and thin areas, as can happen in the *hp*-RB-approach. Instead of two stages of partitioning and then piecewise basis generation, this

approach has a single stage: For a given subdomain a basis generation is started. As soon as it can be predicted that the desired accuracy cannot be met with the currently prescribed maximum basis size, the basis generation is stopped (early stopping greedy), the subdomain is uniformly refined into subdomains and the basis generation is restarted on all child elements. The prediction and early stopping of the greedy procedure is crucial, otherwise $N_{\max}$ basis vectors would have been generated on the coarse element, before one detects that the basis must be discarded and the element must be refined. This prediction therefore is based on an extrapolation procedure, estimating the error decay by the decrease of the error for only a few iterations. For more details we refer to [29].

We want to draw the attention to the conditioning issue: If $\mu^{(i)} \approx \mu^{(j)}$ it can be expected due to continuity that the two snapshots $u(\mu^{(i)}), u(\mu^{(j)})$ are almost linearly dependent. Hence, the corresponding rows/columns of the reduced system matrix $\mathbf{A}_N$ will be almost linearly dependent and hence $\mathbf{A}_N$ possibly be badly conditioned. As seen in Prop. 2.12 orthonormalization of a basis may improve the conditioning of the reduced system. Interestingly, this can be realized via the Gramian matrix, i.e. the matrix $\mathbf{K}_N$ of inner products of the snapshots, and does not involve further expensive high-dimensional operations. For some interesting properties of Gramian matrices, we refer to Exercise 5.14. Using the Gramian matrix, the Gram-Schmidt orthonormalization can then be performed by a Cholesky factorization.

Proposition 2.48 (Orthonormalization of Reduced Basis). *Assume $\Phi_N = \{\varphi_1, \dots, \varphi_N\}$ to be a reduced basis with Gramian matrix denoted $\mathbf{K}_N$. Choose $\mathbf{C} := (\mathbf{L}^T)^{-1} \in \mathbb{R}^{N \times N}$ with $\mathbf{L}$ being a Cholesky-factor of $\mathbf{K}_N = \mathbf{L}\mathbf{L}^T$. We define the transformed basis $\tilde{\Phi}_N := \{\tilde{\varphi}_1, \dots, \tilde{\varphi}_N\}$ by $\tilde{\varphi}_j := \sum_{i=1}^N C_{ij} \varphi_i$. Then $\tilde{\Phi}_N$ is the Gram-Schmidt orthonormalized basis.*

Again the proof is skipped and left as Exercise 5.15. If we are working with a discrete $(P(\mu))$ and the initial basis is given via the coefficient matrix $\Phi_N \in \mathbb{R}^{\mathcal{N} \times N}$ and $\mathbf{K} \in \mathbb{R}^{\mathcal{N} \times \mathcal{N}}$ denotes the full inner-product matrix, then the Gramian matrix is obtained by (27), the matrix $\mathbf{C}$ can be computed as stated in the proposition and the coefficient matrix of the transformed basis is simply obtained by the matrix product $\tilde{\Phi}_N = \Phi_N \mathbf{C}$. Actually, instead of Gram-Schmidt also other transformations are viable, cf. Exercise 5.16.

We briefly comment on basis generation for the primal-dual RB-approach.

Remark 2.49 (Basis Generation for Primal-Dual-Approach). *The a-posteriori error bound in Prop. 2.27 suggests to choose the dimensionality of X_N, X_N^{du} such that $\Delta_u(\mu) \leq \varepsilon_{\text{tol}}$ and $\Delta_u^{\text{du}}(\mu) \leq \varepsilon_{\text{tol}}$ in order to obtain the “squared” effect in the error bound. For construction of X_N, X_N^{du} one could proceed as follows: i) Run independent greedy procedures for the generation of X_N, X_N^{du} by using $(P(\mu))$ and $(P'(\mu))$ as snapshot suppliers, using the same tolerance ε_{tol} and $\Delta_u(\mu)$ and $\Delta_u^{\text{du}}(\mu)$ as error indicator. ii) Run a single greedy procedure based on the output error bound $\Delta'_s(\mu)$ and add snapshots of the solutions of $(P(\mu)), (P'(\mu))$ either in a single space X or in separate spaces X_N, X_N^{du} in each iteration.*

We conclude this section with some experiments which are contained in the script `rb_tutorial.m` in our toolbox *RBmatlab*. In particular, we continue the previous

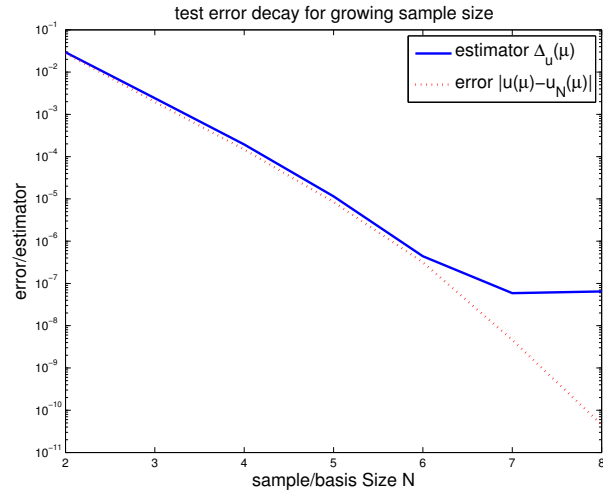


Figure 6: Error convergence for Lagrangian reduced basis with equidistant snapshots.

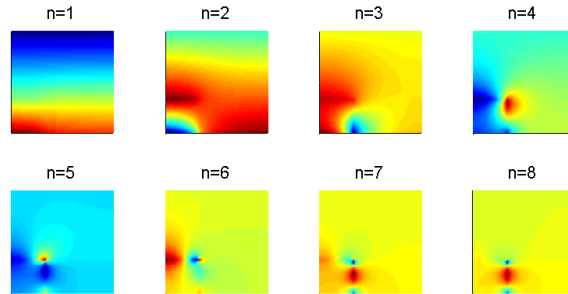


Figure 7: Plot of 8 orthogonal basis functions of an equidistant Lagrangian reduced basis.

experiments on the thermal block model from Sec. 2.1 with a focus on the basis generation. In Fig. 6 we demonstrate the error and error bound convergence when using a Lagrangian reduced basis with equidistantly sampled parameter set. For this example we choose $B_1 = B_2 = 3$ and $\mu = (\mu_1, 1, \dots, 1)$ with $\mu_1 \in [0.5, 2]$. The error and error bound are measured as a maximum over a random test set of parameters S_{test} with $|S_{\text{test}}| = 100$. We observe nice exponential error decay with respect to the sample size N for both the error and the error bound. A typical effect is the flattening or saturation of the error bound at values of about 10^{-7} , when using double values. This is explained by and expected due to numerical accuracy problems, as the error bound is a square root of a residual norm, the accuracy of which is limited by machine precision. The first 8 basis vectors of the corresponding orthonormalized Lagrangian reduced basis are illustrated in Fig. 7. We clearly see the steep variations of the basis functions around the first subblock Ω_1 .



Figure 8: Maximum test error and error bound for greedy reduced basis generation for $B_1 = B_2 = 2$.

Finally, we investigate the results of the greedy procedure. For $\mu_{\min} = 1/\mu_{\max} = 0.5$ and $B_1 = B_2 = 2$ we choose a training set of 1000 random points, $\varepsilon_{\text{tol}} = 10^{-6}$ and the field error estimator as error indicator, i.e. $\Delta(Y, \mu) := \Delta_u(\mu)$. We again measure the quality by determining the maximum error and error bound over a random test set of size 100. The results are plotted in Fig. 8 and nicely confirm the exponential convergence of the greedy procedure for smooth problems. With growing block number $B_1 = B_2 = 2, 3, 4$ the problem is getting more complicated as is indicated in Fig. 9, where the slope of the greedy training estimator curves $(\varepsilon_n)_{n \in N}$ obviously is flattening.

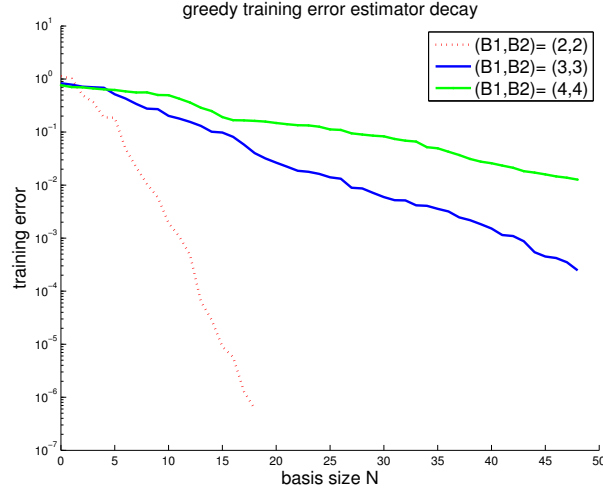


Figure 9: Maximum training error estimator of the greedy procedure for varying block numbers.

2.7 Empirical Interpolation Method

As last approach of this section, we comment on an important general interpolation procedure, the empirical interpolation (EI) method [2, 51], that can be used in RB-methods for several purposes. In particular the procedure can be used in the case, where the given problem does not allow for a parameter separable representation. The EI-method then generates a parameter separable approximation, which can be used as approximate full problem. The notion *empirical* is motivated by the fact that it is based on function snapshots for parametric function interpolation. The EI-method selects a basis, that can be used in addition to the reduced basis, which is therefore called *collateral basis*. The procedure is another instance of a greedy type procedure.

Definition 2.50 (Empirical Interpolation Method, Offline Phase). *Let $\Omega \subset \mathbb{R}^d$ be an open bounded domain, $G \subset C^0(\bar{\Omega})$ a bounded closed set of functions and $\varepsilon_{\text{tol}, \text{EI}}$ a given*

error tolerance. Set $Q_0 := \emptyset, X_0 := \{0\}, T_0 := \emptyset, m = 0$ and define iteratively

$$\begin{aligned}
\text{while } \varepsilon_m &:= \max_{g \in G} \|g - I_m(g)\|_\infty > \varepsilon_{\text{tol}, EI} & (39) \\
g_{m+1} &:= \arg \max_{g \in G} \|g - I_m(g)\|_\infty \\
r_{m+1} &:= g_{m+1} - I_m(g_{m+1}) \\
x_{m+1} &:= \arg \max_{x \in \Omega} |r_{m+1}(x)| \\
T_{m+1} &:= T_m \cup \{x_{m+1}\} \\
q_{m+1} &:= r_{m+1}/r_{m+1}(x_{m+1}) \\
Q_{m+1} &:= Q_m \cup \{q_{m+1}\} \\
X_{m+1} &:= \text{span}(Q_{m+1}) \\
m &\leftarrow m + 1 \\
\text{end while.}
\end{aligned}$$

Here, $I_m : C^0(\bar{\Omega}) \rightarrow X_m$ denotes the interpolation operator with respect to the interpolation points $T_m \subset \Omega$, and the interpolation space X_m , i.e. the unique operator with $I_m(g)(x_{m'}) = g(x_{m'}), m' = 1, \dots, m, g \in G$.

The algorithm produces an interpolation space X_M and points T_M by setting $M := m + 1$ as soon as the loop condition (39) is false.

We give some remarks on practical implementation of the procedure, the assumed function regularity, the collateral basis and the interpolation points.

Remark 2.51 (Practical Implementation). *First, if $|G| = \infty$ then a finite subset $G_{\text{train}} \subset G$ is used instead of G for practical purposes. Similarly, Ω is usually replaced by a finite set of points $\Omega_{\text{train}} \subset \Omega$ to obtain a computable algorithm. Further, instead of the accuracy termination criterion (39), the final dimension M can be specified as input parameter (as long as $M \leq \dim \text{span}(G)$) and the extension loop can be terminated as soon as M is reached.*

Remark 2.52 (Function Regularity). *The formulation assumes continuous functions, which can directly be extended to arbitrary functions, which allow point evaluations. Note, that in this respect L^∞ , as frequently assumed in literature, is not sufficient for a well-defined scheme.*

Remark 2.53 (Collateral Basis). *The sets Q_M are called collateral bases. They are hierarchical in the sense that $Q_M \subset Q_{M+1}$. Further they have the property that $q_m(x_{m'}) = 0$ for $m' < m$.*

Remark 2.54 (Magic Points). *Surprisingly, when choosing a set of polynomials $G = \{x^i\}_{i=1}^n$ on $\bar{\Omega} = [-1, 1]$, the resulting points are such that $\cos^{-1}(T_M)$ is roughly equidistant [27]. This is an interesting fact, as this makes the points have a similar characteristic as the optimal point set for polynomial interpolation, which are the Tscheybysheff points. This motivates the notion “magic points” [51] for the T_M , as the heuristic procedure “magically” produces point sets, which are known to be optimal. Certainly, a fact*

indicating that this is not a real surprise, is that the interpolation points are automatically determined in a greedy fashion to minimize the supremum norm of the resulting interpolation residual. This optimality with respect to the supremum norm is exactly the property of the Tschebysheff points. While theoretically optimal interpolation point sets are only known for simple geometries, the EI is straightforwardly applicable to any arbitrarily shaped domain Ω , rendering it a powerful interpolation technique.

Given the offline data Q_M, T_M of the EI, the online phase, i.e. the actual interpolation, can be easily formulated.

Proposition 2.55 (Empirical Interpolation, Online Phase). *Let a collateral basis Q_M and interpolation points T_M be given from the offline phase of the EI. Then, the matrix*

$$Q_M := (q_j(x_i))_{i,j=1}^M \in \mathbb{R}^{M \times M} \quad (40)$$

is a lower triangular matrix with 1 diagonal, hence regular. Set the vector $g \in C^0(\bar{\Omega})$, $\mathbf{g}_M := (g(x_i))_{i=1}^M \in \mathbb{R}^M$ and let $\alpha_M = (\alpha_i)_{i=1}^M \in \mathbb{R}^M$ be the solution of the linear equation system

$$Q_M \alpha_M = \mathbf{g}_M.$$

Then the interpolation of g is simply

$$I_M(g) = \sum_{i=1}^M \alpha_i q_i. \quad (41)$$

Proof. The fact that the system matrix is lower triangular with 1 diagonal simply follows from the definition. For the interpolation property we directly see that both sides of (41) coincide in the interpolation points $i = 1, \dots, M$:

$$\sum_{j=1}^M \alpha_j q_j(x_i) = \sum_{j=1}^M (Q_M)_{ij} \alpha_j = (Q_M \alpha_M)_i = (\mathbf{g}_M)_i = g(x_i) = I_M(g)(x_i).$$

□

As for other interpolation techniques, the relation to the best approximation is of interest, which is reflected by the Lebesgue constant. For any interpolation operator $I : C^0(\bar{\Omega}) \rightarrow X_M$, the Lebesgue constant is defined as

$$\Lambda_M := \max_{x \in \bar{\Omega}} \sum_{i=1}^M |\xi_i(x)|$$

where ξ_i are nodal basis functions with respect to the interpolation points. Then for any interpolation operator and any $u \in C^0(\bar{\Omega})$ holds

$$\|u - I(u)\|_\infty \leq (1 + \Lambda_M) \inf_{v \in X_M} \|u - v\|_\infty.$$

Now, for the EI the Lebesgue constant can be upper bounded. The following result holds [51].

Proposition 2.56 (Lebesgue Constant Bound). *For the EI-operator I_M the Lebesgue constant can be upper bounded by*

$$\Lambda_M \leq 2^M - 1.$$

and the bound is sharp.

This statement is quite pessimistic, as much better Lebesgue constants can be observed in practice. However, the bound on the Lebesgue constant is sharp, as examples can be constructed, where this bound is attained.

In contrast, an optimistic result is given in the following. Similar as for the (RB) greedy procedure, also for the EI-offline procedure approximation rate statements can be given. For example a central result of [51] is the following.

Proposition 2.57 (A-priori Convergence Rate). *If there exists a sequence of exponentially approximating subspaces, i.e. $Z_1 \subset Z_2 \subset \dots \text{span}(G)$ with $\dim Z_M = M$ for $M \in \mathbb{N}$ and there exist $c > 0, \alpha > \log(4)$ such that*

$$\inf_{v \in Z_M} \|u - v\|_\infty \leq c e^{-\alpha M}, \quad \forall u \in G, M \in \mathbb{N}$$

then the EI offline phase yields almost as good spaces in the sense that

$$\|u - I_M(u)\|_\infty \leq c e^{-(\alpha - \log(4))M}.$$

In addition to such a-priori convergence statements, also a-posteriori error control is possible.

Proposition 2.58 (A-posteriori Error Estimation for EI). *Let $I_M, I_{M'}$ be EI-operators for $M' > M$ with corresponding collateral basis and interpolation points. Set the matrix $\mathbf{Q} := (q_j(x_i))_{i,j=M+1}^{M'} \in \mathbb{R}^{(M'-M+1) \times (M'-M+1)}$. For $g \in C^0(\bar{\Omega})$ let $\mathbf{g}' := (g(x_i) - I_M(g)(x_i))_{i=M+1}^{M'}$ and $\alpha' := (\alpha'_i)_{i=M+1}^{M'} = \mathbf{Q}^{-1} \mathbf{g}'$. If $g \in \text{span}(Q_{M'})$ then the following a-posteriori error bounds hold:*

$$\|g - I_M(g)\|_\infty \leq \Delta_{\text{EI},\infty}(g) := \|\alpha'\|_1 = \sum_{i=M+1}^{M'} |\alpha'_i|, \quad (42)$$

$$\|g - I_M(g)\|_{L^2} \leq \Delta_{\text{EI},2}(g) := \sqrt{(\alpha')^T \mathbf{K}_Q \alpha'} \quad (43)$$

where

$$\mathbf{K}_Q := \left(\int_{\Omega} q_i q_j \right)_{i,j=M+1}^{M'}.$$

The proof is left as a simple exercise using the nestedness of the interpolation matrices Q_M and $Q_{M'}$ and the definitions.

In this bound, a certain exactness for a higher EI-index $M' > M$ is assumed for exact certification. This assumption can be widely found, e.g. [24, 73]. An alternative, which requires a certain smoothness knowledge instead of this exactness assumption, has been presented in [18].

A useful property is the following, which states conservation of linear operations.

Proposition 2.59 (Conservation Property). *Let $f \in (C^0(\bar{\Omega}))'$ be a linear functional on continuous functions and $f(g) = 0$ for all $g \in G$. Then we also have*

$$f(I_M(g)) = 0, \quad \forall g \in G.$$

Proof. By linearity we immediately have $f(g)$ for all $g \in \text{span}(G)$. As $Q_M \subset \text{span}(G)$ by construction, the claim follows by linearity of the interpolation operator. $\square$

An intuitive example of such functionals are conservation of zeros/roots: If f is a point evaluation in $\bar{x}$ and $g(\bar{x}) = 0$ for all $g \in G$, then also $I_M(g)(\bar{x}) = 0$. This can for instance be used in the conservation of zero entries in the interpolation of sparse vectors or (vectorized) sparse matrices [73].

Another example could be the interpolation of zero-mean functions: If $\int_{\Omega} g = 0$ for all $g \in G$, then also $\int_{\Omega} I_M(g) = 0$. This can be beneficial in the interpolation of conservative operators [17]. Similarly, one could imagine interpolation of divergence free velocity fields in fluid dynamics, etc.

Now we will establish the connection to the previous sections, i.e. apply the EI in a parametric context and formulate an EI-RB-scheme. In particular we provide EI-approximations for problems with non-separable data functions.

Definition 2.60 (EI for Parametric Functionals). *Let $X \subset L^2(\Omega)$ and $f \in X'$ be a parametric continuous linear form of the type*

$$f(v; \mu) = \int_{\Omega} g(x; \mu) v(x) dx \quad (44)$$

for $g(\cdot; \mu) \in C^0(\bar{\Omega})$. In general g and f may be non parameter separable. Then set $G := \{g(\cdot; \mu) | \mu \in \mathcal{P}\}$ and compute the EI offline data $Q_M = \{q_i\}_{i=1}^M$ and $T_M = \{x_i\}_{i=1}^M$ according to Def. 2.50. Then we obtain parameter-separable approximations $\tilde{g}, \tilde{f}$ for g and f by

$$\tilde{g}(\cdot; \mu) := I_M(g(\cdot; \mu)) = \sum_{m=1}^M \theta_m(\mu) \tilde{g}_m(\cdot) \quad (45)$$

$$\tilde{f}(v; \mu) := \int_{\Omega} \tilde{g}(x; \mu) v(x) dx = \sum_{m=1}^M \theta_m(\mu) \tilde{f}_m \quad (46)$$

with components $\tilde{g}_m := q_m$, $\tilde{f}_m := \int_{\Omega} \tilde{g}_m v$ and coefficient function vector

$$(\theta_1(\mu), \dots, \theta_M(\mu))^T := Q_M^{-1} \mathbf{g}(\mu)$$

using the local evaluation vector $\mathbf{g}(\mu) := (g(x_1; \mu), \dots, g(x_M; \mu))^T$ and interpolation matrix Q_M as in (40).

Similar parameter separable approximation can be obtained for non-parameter separable bilinear forms, e.g. $a(u, v; \mu) := \int_{\Omega} \kappa(x; \mu) \nabla u(x) \cdot \nabla v(x) dx$ by EI for κ .

Hereby, a non-parameter separable problem of type (P) with solution u can be approximated by EI of the data functions resulting in an approximate variational form ($\tilde{P}$) with solution $\tilde{u}$.

For simplicity, in the following we ignore the output functional.

Proposition 2.61 (EI-approximated Full Problem ($\tilde{P}(\mu)$)). *Assume that $a(\cdot, \cdot; \mu)$ is a continuous bilinear form, uniformly coercive in μ and $f(\cdot; \mu)$ is uniformly bounded with respect to μ . Assume that the EI-approximations are sufficiently accurate in the sense that we have $\varepsilon_a, \varepsilon_f \in \mathbb{R}$ with $\varepsilon_a < \tilde{\alpha}$ such that for all $u, v \in X, \mu \in \mathcal{P}$*

$$|a(u, v; \mu) - \tilde{a}(u, v; \mu)| \leq \varepsilon_a \|u\| \|v\|, \quad |f(v; \mu) - \tilde{f}(v; \mu)| \leq \varepsilon_f \|v\|. \quad (47)$$

Then the forms $\tilde{a}, \tilde{f}$ are continuous and $\tilde{a}$ is coercive with coercivity lower bound $\tilde{\alpha} := \alpha - \varepsilon_a > 0$, hence the following problem has a unique solution $\tilde{u}(\mu) \in X$

$$\tilde{a}(\tilde{u}, v; \mu) = \tilde{f}(v; \mu), \quad v \in X \quad (48)$$

which satisfies $\|\tilde{u}\| \leq \|\tilde{f}\|_{X'} / \tilde{\alpha}$.

Proof. Continuity of the approximated forms simply follows by

$$|\tilde{f}(v)| \leq |\tilde{f}(v) - f(v)| + |f(v)| \leq \varepsilon_f \|v\| + \|f\|_{X'} \|v\| = (\varepsilon_f + \|f\|_{X'}) \|v\|$$

and similar for $\tilde{a}$. For the coercivity we obtain

$$\begin{aligned} \frac{\tilde{a}(u, u)}{\|u\|^2} &= \frac{a(u, u) - (a(u, u) - \tilde{a}(u, u))}{\|u\|^2} \geq \frac{a(u, u)}{\|u\|^2} - \frac{|a(u, u) - \tilde{a}(u, u)|}{\|u\|^2} \\ &\geq \alpha - \varepsilon_a = \tilde{\alpha} > 0. \end{aligned}$$

The remaining statements follow by the Lax-Milgram theorem. $\square$

This statement is valid for any approximation procedure. Specifically for the EI and the previous example for f in (44), it is easy to see that the ε_f can be related to the error bounds (assuming their validity):

$$|f(v) - \tilde{f}(v)| = \left| \int_{\Omega} (g - \tilde{g})v \right| \leq \|g - \tilde{g}\|_{L^2} \|v\|_{L^2} \leq \Delta_{\text{EI},2}(g(\cdot; \mu)) \|v\|_{H^1}.$$

Hence, choosing $\varepsilon_f \geq \sup_{\mu \in \mathcal{P}} \Delta_{\text{EI},2}(g(\cdot; \mu))$ will guarantee the validity on the f -approximation assumption in (47). Similarly, for the bilinear form $a(u, v) = \int_{\Omega} \kappa(\cdot; \mu) \nabla u \cdot \nabla v$ we can satisfy the assumption on a in (47) by choosing $\varepsilon_a \geq \sup_{\mu \in \mathcal{P}} \Delta_{\text{EI},\infty}(\kappa(\cdot; \mu))$. So, ε_a can be made small by ensuring that the $\Delta_{\text{EI},\infty}(\kappa(\cdot; \mu))$ are small, i.e. the EI is sufficiently accurate. Now, the RB-machinery of the previous sections can be applied to generate an approximate well-posed reduced problem ($\tilde{P}_N$) with solution $\tilde{u}_N$ and the error $\tilde{u} - \tilde{u}_N$ can be quantified by the presented estimators. However, for controlling the complete error $u - \tilde{u}_N$ the interpolation error needs to be additionally estimated. This can be obtained by a disturbance argument.

Proposition 2.62 (A-posteriori Error Estimator for EI-RB-Approximation). *Let $u(\mu) \in X$ be the solution of the non parameter separable problem (P) and $\tilde{u}_N(\mu)$ be the RB-approximation of the EI-approximated system ($\tilde{P}(\mu)$). Then we have the error bound*

$$\|u - \tilde{u}_N\| \leq \Delta_{\text{EI}}(\mu) + \Delta_{\tilde{u}}(\mu) \quad (49)$$

where $\Delta_{\tilde{u}}(\mu)$ denotes the standard RB error bound for the error $\tilde{u} - \tilde{u}_N$ analogous to (8) and Δ_{EI} is an appropriate empirical interpolation error contribution

$$\Delta_{\text{EI}}(\mu) := \frac{1}{\alpha} \varepsilon_f + \frac{\|\tilde{f}\|_{X'}}{\alpha(\alpha - \varepsilon_a)} \varepsilon_a. \quad (50)$$

Proof. We first note by definitions of (P) and ($\tilde{P}$)

$$a(u - \tilde{u}, v) = a(u, v) - a(\tilde{u}, v) = f(v) - \tilde{f}(v) + \tilde{a}(\tilde{u}, v) - a(\tilde{u}, v) =: \tilde{r}(v),$$

where the second equality follows from adding $0 = \tilde{f}(v) - \tilde{a}(\tilde{u}, v) = 0$. Using the approximation property yields

$$\|\tilde{r}\|_{X'} \leq \varepsilon_f + \varepsilon_a \|\tilde{u}\|.$$

Then, Lax-Milgram together with the bound for $\tilde{u}$ from Prop. 2.61 allows to conclude

$$\|u - \tilde{u}\| \leq \frac{1}{\alpha} \|\tilde{r}\| \leq \frac{1}{\alpha} \varepsilon_f + \frac{\|\tilde{f}\|_{X'}}{\alpha(\alpha - \varepsilon_a)} \varepsilon_a = \Delta_{\text{EI}}(\mu). \quad (51)$$

The overall bound (49) then follows by the triangle inequality. $\square$

3 Instationary Problems

In this section, we aim at extending the methodology to time-dependent problems. Historically, time-dependent problems have even been the motivation for reduced basis modelling [1, 4, 41, 58], in particular in the context of fluid flow and for the purpose of understanding complexity in turbulence. However, these techniques were not certified by error estimation and not subject to offline/online decomposition, etc. as presented for the stationary case. The first publication known to us dealing with *certified* RB-method for instationary problems can be found in [24, 26], where parabolic problems were considered. A central result of that work, a space-time energy norm error estimator, will be given in a reformulated fashion in this section. Otherwise, the formulations and techniques of this section are now based on our previous work [31] and comprise also new results by following the pattern of the previous section. In particular we will proceed in parallel to the stationary case and sequentially prove similar results with the same notions.

Instead of giving a variational formulation, we consider an alternative operator-based formulation in the current section. This will in particular allow an RB-approach

for finite difference or finite volume discretizations of hyperbolic equations, which are usually not motivated by a variational formulation. However, note that all of the following could as well be formulated in a variational fashion.

We admit that by the focus and choice of the presentation in this section we are clearly biased to our own previous work, as most other articles on RB-methods for instationary problems with variational discretizations use corresponding weak forms of the PDE. For such formulations and RB-schemes, we refer to [20, 19, 25, 24, 26, 44, 47, 67].

3.1 Model Problem

As model problem we consider a linear advection-diffusion problem on a rectangular domain $\Omega = (0, 2) \times (0, 1)$ with end time $T = 1$, i.e. for given μ find $u(x, t; \mu)$ as solution of

$$\partial_t u(\mu) + \nabla \cdot (\mathbf{v}(\mu)u(\mu) - d(\mu)\nabla u(\mu)) = 0 \quad \text{in } \Omega \times (0, T)$$

with suitable initial condition $u(x, 0; \mu) = u_0(x; \mu)$ and Dirichlet boundary conditions. The initial and time-variant inhomogeneous Dirichlet boundary values $g_D(x, t)$ are based on a nonnegative radial basis function linearly decaying over time with center on the top edge at $x_1 = 1/2$. The velocity field is chosen as a superposition of two divergence free parabolic velocity fields in x_1 and x_2 -direction with weighting factor 1 and μ_1 , i.e.

$$\mathbf{v}(x; \mu) = \left(\mu_1 \frac{5}{2}(1 - x_2^2), -\frac{1}{2}(4 - x_1^2) \right)^T.$$

The diffusivity is chosen as $d(\mu) := 0.03 \cdot \mu_2$ resulting in a parametrization of the velocity and diffusivity by $(\mu_1, \mu_2) \in [0, 1]^2$. Note, that the above problem is changing type with parameter, i.e. for $\mu_2 = 0$ we obtain a hyperbolic problem, while for $\mu_2 > 0$ it is parabolic. Some solution snapshots over time (using a FV-discretization, details will be reported in the experiments at the end of Sec. 3.5) are presented in Fig. 10, each column representing a time evolution of a different parameter.

3.2 Full Problem

The above problem is an example of a general linear evolution problem of the type

$$\begin{aligned} \partial_t u - \mathcal{L}(u; \mu) &= q(\mu) \quad \text{in } \Omega \times (0, T) \\ u(0) &= u_0(\mu) \quad \text{in } \Omega \end{aligned}$$

With $\Omega \subset \mathbb{R}^d$ the spatial domain, $[0, T]$ the time-interval with final time $T > 0$, $\mathcal{L}$ denotes a linear spatial differential operator, q an inhomogeneity and u_0 the initial values.

The full problem will be based on a time-discrete formulation based on $K \in \mathbb{N}$ steps in time, stepsize $\Delta t := T/K$ and time instants $t^k := k\Delta t, k = 0, \dots, K$. For notational convenience, we assume constant Δt , though varying time-step widths can easily be incorporated in the following.

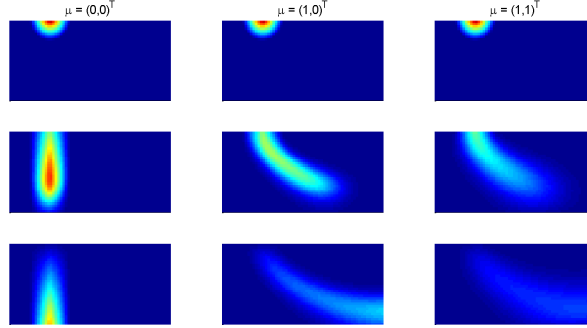


Figure 10: Illustration of snapshots of the advection diffusion example: initial data (top) and time evolution at time $t = 0.5$ (middle row) and $t = 1$ (bottom) for parameter vectors $\mu = (0,0)^T, (1,0)^T, (1,1)^T$ from left to right.

We again assume that X is a Hilbert space with inner product $\langle \cdot, \cdot \rangle$ and norm $\|\cdot\|$ and we seek the solution variable $u^k(\mu) \in X$ with $u^k(x; \mu) \approx u(x, t^k; \mu)$ for $k = 0, \dots, K$. We assume a general explicit/implicit time-discretization with $\mathcal{L}_I^k(\mu), \mathcal{L}_E^k(\mu) : X \rightarrow X$ linear continuous operators and $b^k(\mu) \in X$. Note, that $\mathcal{L}_I^k, \mathcal{L}_E^k, b^k$ will typically depend on Δt , hence these operators reflect both the time-discretization as well as the space discretization. For simplicity we assume $u^0 \in X$, otherwise an initial data projection needs to be included. Then, the problem for the discrete solution can be formulated as a time-marching evolution scheme.

Definition 3.1 (Full Evolution Problem ($E(\mu)$)). *For $\mu \in \mathcal{P}$ find a sequence of solutions $\{u^k(\mu)\}_{k=0}^K \subset X$ by starting with $u^0(\mu) \in X$ and iteratively solving the following operator equations for $u^{k+1}(\mu)$*

$$\mathcal{L}_I^k(\mu)u^{k+1}(\mu) = \mathcal{L}_E^k(\mu)u^k(\mu) + b^k(\mu), \quad k = 0, \dots, K-1.$$

We omit output estimation here, and give comments on this in Rem. 3.13. We again formulate some requirements for well-posedness.

Definition 3.2 (Uniform Continuity and Coercivity). *The parametric operators are assumed to be continuous with continuity constants $\gamma_I^k(\mu) := \|\mathcal{L}_I^k(\mu)\|, \gamma_E^k(\mu) := \|\mathcal{L}_E^k(\mu)\|$. The continuity is assumed to be uniform with respect to μ and t in the sense that for some $\bar{\gamma}_I, \bar{\gamma}_E < \infty$ holds $\gamma_I^k(\mu) \leq \bar{\gamma}_I, \gamma_E^k(\mu) \leq \bar{\gamma}_E$ for all μ and k . Further, $\mathcal{L}_I^k$ is assumed to be coercive, i.e. there exists a constant*

$$\alpha_I^k(\mu) := \inf_{v \in X \setminus \{0\}} \frac{\langle \mathcal{L}_I^k(\mu)v, v \rangle}{\|v\|^2} > 0$$

and the coercivity is uniform with respect to μ and t in the sense that for some $\bar{\alpha}_I > 0$ holds $\alpha_I^k(\mu) \geq \bar{\alpha}_I$ for all μ and k . Similarly, for $b^k(\mu)$ we assume uniform boundedness by $\|b^k(\mu)\| \leq \bar{\gamma}_b$ for suitable $\bar{\gamma}_b$.

Under these assumptions, one obtains well-posedness and stability of the problem $(E(\mu))$.

Proposition 3.3 (Well-posedness and Stability of $(E(\mu))$). *The solution trajectory $\{u^k(\mu)\}_{k=0}^K$ of $(E(\mu))$ is well-defined and bounded by*

$$\|u^k(\mu)\| \leq \|u^0\| \left(\frac{\bar{\gamma}_E}{\bar{\alpha}_I} \right)^k + \frac{\bar{\gamma}_b}{\bar{\alpha}_I} \left(\sum_{i=0}^{k-1} \left(\frac{\bar{\gamma}_E}{\bar{\alpha}_I} \right)^i \right). \quad (52)$$

Proof. Well-definedness of the solution in iteration k follows by Lax-Milgram, uniform continuity/coercivity and gives the bound

$$\|u^{k+1}\| \leq \frac{1}{\bar{\alpha}_I} \left(\bar{\gamma}_E \|u^k\| + \bar{\gamma}_b \right).$$

The bound (52) then easily follows by induction. $\square$

The constants $\bar{\gamma}_E, \bar{\gamma}_b$ and $\bar{\alpha}_I$, and herewith the constant on the right hand side of (52) depend on Δt . Hence the behavior for $\Delta t \rightarrow 0$ is of interest. One can show that the solution does not diverge with decreasing Δt , under some conditions on the continuity and coercivity constant. We leave the proof of the following as Exercise 5.17.

Proposition 3.4 (Uniform Boundedness with respect to Δt). *Let $\bar{\gamma}_E \leq 1$, $\bar{\alpha}_I = 1 + \alpha \Delta t$ and $\bar{\gamma}_b = C \Delta t$ with α, C independent of Δt . Then*

$$\lim_{K \rightarrow \infty} \|u^K\| \leq e^{-\alpha T} \|u^0\| + \tilde{C} T$$

with explicitly computable, Δt -independent constant $\tilde{C}$.

Note, that this statement is a very coarse qualitative statement. Certainly stronger results such as convergence of $u^K(\mu)$ (and for any other t) are usually expected (and provided) by reasonable discretizations, i.e. instantiations of $(E(\mu))$.

Again, we assume parameter-separability analogous to Def. 2.5 for later efficient offline/online-decomposition. In contrast to the stationary case, we assume that the time dependency also is encoded in the coefficient functions such that the components are parameter- and time-independent.

Definition 3.5 (Parameter-Separability). *We assume the operators $\mathcal{L}_I^k, \mathcal{L}_E^k, b^k$ to be parameter-separable, i.e. there exist coefficient functions $\theta_{I,q}^k : \mathcal{P} \rightarrow \mathbb{R}$ and parameter-independent continuous linear operators $\mathcal{L}_{I,q} : X \rightarrow X$ for $q = 1, \dots, Q_I$ such that*

$$\mathcal{L}_I^k(\mu) = \sum_{q=1}^{Q_I} \theta_{I,q}^k(\mu) \mathcal{L}_{I,q}$$

and similar definitions for $\mathcal{L}_E^k, b^k$ and u^0 with corresponding coefficient functions $\theta_{E,q}^k(\mu), \theta_{b,q}^k(\mu), \theta_{u^0,q}(\mu)$, components $\mathcal{L}_{E,q}, b_q, u_q^0$ and number of components Q_E, Q_b, Q_{u^0} .

Assuming Lipschitz-continuity of the coefficient functions with respect to μ , one can derive Lipschitz-continuity of the solution similar to Prop. 2.7. Also, sensitivity equations for instationary problems can be obtained similar to Prop. 2.8, which is made use of in the context of parameter optimization, cf. [15].

Several examples fit into the framework of the above problem $(E(\mu))$.

Example 7 (Finite Element Formulation for a Diffusion Problem). *For an implicit finite element discretization of a diffusion problem with homogeneous Dirichlet boundary data we can choose $X := \text{span}\{\phi_i\}_{i=1}^{\mathcal{N}} \subset H_0^1(\Omega)$ as the space of piecewise linear functions on the grid $\mathcal{T}$ assigned with the $H_0^1(\Omega)$ inner product. The variational time marching form of parabolic problems*

$$m(u^{k+1}, v) + \Delta t a(u^{k+1}, v) = m(u^k, v), \quad v \in X$$

with $m(u, v) := \int_{\Omega} u v, a(u, v) := \int_{\Omega} d \nabla u \cdot \nabla v$ and $L^2(\Omega)$ orthogonal projection of the initial data $u^0 := P_X u_0$ then can directly be transferred to the operator formulation by defining $b^k := 0$ and the operators implicitly via

$$\langle \mathcal{L}_E^k u, v \rangle := m(u, v), \quad \langle \mathcal{L}_I^k u, v \rangle := m(u, v) + \Delta t a(u, v), \quad u, v \in X.$$

Parameter-separability of the data functions will result in parameter-separable operators. However, note that this Galerkin formulation does not allow (large) advection terms or vanishing diffusion unless accepting possible instability, as the implicit spatial discretization operator may become non-coercive.

Example 8 (Finite Volume Discretization for Advection Problem). *Given a triangulation $\mathcal{T} = \{T_i\}_{i=1}^{\mathcal{N}}$ of $\Omega \subset \mathbb{R}^d$ one can choose the discrete basis functions $\psi_i := \chi_{T_i}, i = 1, \dots, \mathcal{N}$ of characteristic functions of the grid elements. Then we define $X := \text{span}\{\psi_i\}_{i=1}^{\mathcal{N}} \subset L^2(\Omega)$ a Finite Volume space of piecewise constant functions with the $L^2(\Omega)$ inner product. An explicit Euler forward time-discretization of an advection problem with Dirichlet boundary data b_{dir} and velocity field $\mathbf{v}$ can then easily be formulated. For this we first define $u^0 = P_X u_0 \in X$ as L^2 -projection of given initial data u_0 to piecewise constant functions and define $\mathcal{L}_I^k := Id$ as the identity on X . Using for example a Lax-Friedrichs numerical flux with parameter $\lambda > 0$ [49], we obtain $b^k = \sum_i b_i^k \psi_i \in X$ with*

$$b_i^k := -\frac{\Delta t}{|T_i|} \sum_{j \in N_{\text{dir}}(i)} \frac{|e_{ij}|}{2} [\mathbf{v}(\mathbf{c}_{ij}) \cdot \mathbf{n}_{ij} - \lambda^{-1}] b_{\text{dir}}(\mathbf{c}_{ij}),$$

where $N_{\text{dir}}(i)$ is the set of indices enumerating the Dirichlet boundary edges of T_i , $|e_{ij}|$ is the length and $\mathbf{c}_{ij}$ the centroid of the corresponding edge and $|T_i|$ indicates the volume of element T_i . The operator $\mathcal{L}_E^k$ is specified via its operation on a vector of unknowns, i.e. for all $\mathbf{w} = \sum_i w_i \psi_i$ and $\mathbf{w}' := \mathcal{L}_E^k \mathbf{w} = \sum_i w'_i \psi_i$ with vectors of unknowns $\mathbf{w}, \mathbf{w}' \in \mathbb{R}^{\mathcal{N}}$

we assume $\mathbf{w}' = \hat{\mathbf{L}}_E^k \mathbf{w}$ with matrix $\hat{\mathbf{L}}_E^k \in \mathbb{R}^{\mathcal{N} \times \mathcal{N}}$ defined by entries

$$(\hat{\mathbf{L}}_E^k)_{i,l} := \begin{cases} 1 - \frac{\Delta t}{|T_i|} \sum_{j \in N_{\text{int}}(i) \cup N_{\text{dir}}(i)} \frac{|e_{ij}|}{2} [\mathbf{v}(\mathbf{c}_{ij}) \cdot \mathbf{n}_{ij} + \lambda^{-1}] & \text{for } l = i \\ -\frac{\Delta t}{|T_i|} \frac{|e_{il}|}{2} [\mathbf{v}(\mathbf{c}_{il}) \cdot \mathbf{n}_{il} + \lambda^{-1}] & \text{for } l \in N_{\text{int}}(i) \\ 0 & \text{otherwise} \end{cases}$$

where $N_{\text{int}}(i)$ are the neighboring element indices of T_i . Note that this discretization scheme requires a CFL condition, i.e. sufficiently small Δt in order to guarantee a stable scheme. Similarly, diffusion terms or additional Neumann boundary values can be discretized as explicitly specified in [31]. Our additional requirements are satisfied under suitable assumptions on the data functions: Again, parameter-separability and uniform continuity of the data functions result in parameter-separability and uniform continuity of the operators. The implicit operator (the identity) clearly is uniformly coercive independent of the data functions. In particular, this example gives a discretization and hence results in an RB-method for a hyperbolic problem. Also, it can be shown that the assumptions of Prop. 3.4 are satisfied.

Example 9 (Finite Differences). Similar to the Finite Volume example, also Finite Difference discretizations can be treated with the current evolution and RB-formulation. Given $\Omega_h = \{x_i\}_{i=1}^{\mathcal{N}} \subset \Omega$ as discrete set of grid points used for the finite difference discretization, we define $X := \text{span}\{\delta_{x_i}\}_{i=1}^{\mathcal{N}}$ as the set of indicator functions $\delta_{x_i} : \Omega_h \rightarrow \mathbb{R}$ of the discrete point set, and assign it with a discrete inner product, e.g. an approximation of a L^2 inner product $\langle u, v \rangle := \sum_{i=1}^{\mathcal{N}} w_i u(x_i) v(x_i)$ for $u, v \in X$ with weights $w_i \in \mathbb{R}$. Again, parameter-separability of the data functions will result in parameter-separable finite difference operators. Uniform boundedness of the data in combination with continuous evaluation functionals results in uniformly bounded operators.

To conclude this section, we summarize that $(E(\mu))$ captures quite general PDEs (hyperbolic, parabolic), which have a time-derivative of first order and arbitrary spatial differential operator. Further, different spatial discretization techniques are allowed: FV, FD, FE or DG schemes, etc. Different time-discretizations are allowed: Euler forward/backward, or Crank-Nicolson. Also operator splitting is allowed: different parts of the continuous differential operator $\mathcal{L}$ may be discretized implicitly, others explicitly.

3.3 RB-Approach

We again assume the availability of an RB-space X_N with reduced basis $\Phi_N = \{\varphi_1, \dots, \varphi_N\}$, $N \in \mathbb{N}$. We give a procedure for basis generation in Sec. 3.5.

Definition 3.6 (RB-Problem $(E_N(\mu))$). For $\mu \in \mathcal{P}$ find a sequence of solutions $\{u_N^k(\mu)\} \subset X_N$ by starting with $u_N^0(\mu) := P_{X_N} u^0(\mu)$ and iteratively solving

$$\mathcal{L}_{I,N}^k(\mu) u_N^{k+1}(\mu) = \mathcal{L}_{E,N}^k(\mu) u_N^k(\mu) + b_N^k(\mu), \quad k = 0, \dots, K-1,$$

with reduced operators and reduced inhomogeneity

$$\mathcal{L}_{N,I}^k(\mu) = P_{X_N} \circ \mathcal{L}_I^k(\mu), \quad \mathcal{L}_{N,E}^k(\mu) = P_{X_N} \circ \mathcal{L}_E^k(\mu), \quad b_N^k(\mu) = P_{X_N} b^k(\mu),$$

where $P_{X_N} : X \rightarrow X_N$ denotes the orthogonal projection with respect to $\langle \cdot, \cdot \rangle$.

The well-posedness and stability of $\{u_N^k\}$ follow exactly identical to Prop. 3.3. We again can state a simple consistency property, which is very useful for validating program code as commented in Rem. 2.17.

Proposition 3.7 (Reproduction of Solutions). *If for some $\mu \in \mathcal{P}$ we have $\{u^k(\mu)\}_{k=0}^K \subset X_N$ then $u_N^k(\mu) = u^k(\mu)$ for $k = 0, \dots, K$.*

Proof. The statement follows by induction. For $k = 0$ we have $u^0 \in X_N$ and $P_{X_N}|_{X_N} = Id$, therefore $u_N^0 = P_{X_N} u^0 = u^0$. For the induction step assume that $u^k = u_N^k$. With $(E_N(\mu))$ we obtain

$$0 = \mathcal{L}_{I,N}^k u_N^{k+1} - \mathcal{L}_{E,N}^k u_N^k - b_N^k = P_{X_N} (\mathcal{L}_I^k u_N^{k+1} - \mathcal{L}_E^k u_N^k - b^k). \quad (53)$$

Using $(E(\mu))$ and $u^k = u_N^k$ we verify $-\mathcal{L}_E^k u_N^k - b^k = -\mathcal{L}_I^k u^{k+1}$ and (53) reduces to $P_{X_N} (\mathcal{L}_I^k (u_N^{k+1} - u^{k+1})) = 0$. This means $\mathcal{L}_I^k e^{k+1} \perp X_N$ using the abbreviation $e^{k+1} := u^{k+1} - u_N^{k+1}$. But by the assumption $u^{k+1} \in X_N$ we also have $e^{k+1} \in X_N$. Hence uniform coercivity implies

$$0 = \langle \mathcal{L}_I^k e^{k+1}, e^{k+1} \rangle \geq \bar{\alpha}_I \|e^{k+1}\|^2$$

proving $e^{k+1} = 0$ and, hence, $u^{k+1} = u_N^{k+1}$. $\square$

Proposition 3.8 (Error-Residual Relation). *For $\mu \in \mathcal{P}$ we define the residual $\mathcal{R}^k(\mu) \in X$ via*

$$\mathcal{R}^k(\mu) := \frac{1}{\Delta t} (\mathcal{L}_E^k(\mu) u_N^k(\mu) - \mathcal{L}_I^k(\mu) u_N^{k+1}(\mu) + b^k(\mu)), \quad k = 0, \dots, K-1. \quad (54)$$

Then, the error $e^k(\mu) := u^k(\mu) - u_N^k(\mu) \in X$ satisfies the evolution problem

$$\mathcal{L}_I^k(\mu) e^{k+1}(\mu) = \mathcal{L}_E^k(\mu) e^k(\mu) + \mathcal{R}^k(\mu) \Delta t, \quad k = 0, \dots, K-1. \quad (55)$$

Proof. Using $(E(\mu))$ and $(E_N(\mu))$ yields

$$\begin{aligned} \mathcal{L}_I^k e^{k+1} &= \mathcal{L}_I^k u^{k+1} - \mathcal{L}_I^k u_N^{k+1} \\ &= \mathcal{L}_E^k u^k + b^k - \mathcal{L}_E^k u_N^k + \mathcal{L}_E^k u_N^k - \mathcal{L}_I^k u_N^{k+1} = \mathcal{L}_E^k e^k + \Delta t \mathcal{R}^k. \end{aligned}$$

$\square$

The following a-posteriori error bound simply follows by applying the a-priori bounding technique of Prop. 3.3 to the error-evolution of Prop. 3.8.

Proposition 3.9 (A-posteriori Error Bound X -norm). *Let $\gamma_{\text{UB}}(\mu)$ and $\alpha_{\text{LB}}(\mu)$ be computable upper/lower bounds satisfying*

$$\gamma_E^k(\mu) \leq \gamma_{\text{UB}}(\mu) \leq \bar{\gamma}_E, \quad \alpha_I^k(\mu) \geq \alpha_{\text{LB}}(\mu) \geq \bar{\alpha}_I, \quad \mu \in \mathcal{P}, k = 0, \dots, K.$$

Then, the RB-error can be bounded by

$$\begin{aligned} \|u^k(\mu) - u_N^k(\mu)\| &\leq \Delta_u^k(\mu) \text{ with} \\ \Delta_u^k(\mu) &:= \|e^0\| \left(\frac{\gamma_{\text{UB}}(\mu)}{\alpha_{\text{LB}}(\mu)} \right)^k + \sum_{i=0}^{k-1} \left(\frac{\gamma_{\text{UB}}(\mu)}{\alpha_{\text{LB}}(\mu)} \right)^{k-i-1} \frac{\Delta t}{\alpha_{\text{LB}}(\mu)} \|\mathcal{R}^i\|. \end{aligned}$$

Proof. The error-residual relation (55) and the Lax-Milgram theorem imply the recursion

$$\|e^{k+1}\| \leq \frac{\gamma_{\text{UB}}(\mu)}{\alpha_{\text{LB}}(\mu)} \|e^k\| + \frac{\Delta t}{\alpha_{\text{LB}}(\mu)} \|\mathcal{R}^k\|.$$

Then, the bound follows by induction. $\square$

The bound can be simplified by ensuring that $u_q^0 \in X_N$: Then by linear combination $u^0(\mu) \in X_N$ and with reproduction of solutions, see Prop. 3.7, we get $\|e^0\| = 0$, consequently the corresponding term in the error bound vanishes. We will return to this in Rem. 3.24.

Note, that there is no simple way to obtain effectivity bounds for this error estimator. Numerically, effectivities of this error bound may be large. Nevertheless, the error bound is rigorous, thus if the error bound is small, the true error is ensured also to be small, usually even some orders of magnitude better. Additionally, the error bound predicts zero error a-posteriori:

Proposition 3.10 (Vanishing Error Bound). *If $u^k(\mu) = u_N^k(\mu)$ for all $k = 0, \dots, K$ then $\Delta_u^k(\mu) = 0, k = 0, \dots, K$.*

Proof. If $u^k = u_N^k, k = 0, \dots, K$ we conclude with $(E(\mu))$ and the residual definition (54) that $\mathcal{R}^k = 0$ and hence $\Delta_u^k(\mu) = 0$. $\square$

The above general estimators can be improved, if more structure or knowledge about the problem is available. The following is a result from [24, 26] and applies to implicit discretizations of symmetric parabolic problems, cf. Example 7. For this, we additionally assume

$$\mathcal{L}_E^k(\mu) = \mathcal{L}_m(\mu), \quad \mathcal{L}_I^k(\mu) = \mathcal{L}_m(\mu) + \Delta t \mathcal{L}_a(\mu), \quad (56)$$

for $k = 0, \dots, K$ and $\mathcal{L}_a, \mathcal{L}_m : X \rightarrow X$ being continuous linear operators independent of t and Δt . Here $\mathcal{L}_m$ can correspond to a general mass term and $\mathcal{L}_a$ can represent a stiffness term. Correspondingly, we additionally assume symmetry

$$\langle \mathcal{L}_m(\mu)u, v \rangle = \langle u, \mathcal{L}_m(\mu)v \rangle, \quad \langle \mathcal{L}_a(\mu)u, v \rangle = \langle u, \mathcal{L}_a(\mu)v \rangle, \quad u, v \in X, \mu \in \mathcal{P}, \quad (57)$$

positive definiteness of $\mathcal{L}_m$ and coercivity of $\mathcal{L}_a$:

$$\langle \mathcal{L}_m(\mu)u, u \rangle > 0, u \neq 0, \quad \alpha(\mu) := \inf_{u \neq 0} \frac{\langle \mathcal{L}_a(\mu)u, u \rangle}{\|u\|^2} > 0, \mu \in \mathcal{P}. \quad (58)$$

Then, $\mathcal{L}_m, \mathcal{L}_a$ induce scalar products and we can define the following μ -dependent space-time energy norm for $u = (u^k)_{k=1}^K \in X^K$

$$\|u\|_\mu := \left(\langle \mathcal{L}_m(\mu)u^K, u^K \rangle + \Delta t \sum_{k=1}^K \langle \mathcal{L}_a(\mu)u^k, u^k \rangle \right)^{1/2}$$

Then the following error bound can be proven.

Proposition 3.11 (A-posteriori Error Bound, Space-Time Energy Norm). *Under the assumptions (56) – (58) we have for the solutions $u(\mu) := (u^k(\mu))_{k=1}^K$, $u_N(\mu) := (u_N^k(\mu))_{k=1}^K$ of $(E(\mu))$ and $(E_N(\mu))$*

$$\|u(\mu) - u_N(\mu)\|_\mu \leq \Delta_\mu^{en}(\mu) \text{ with}$$

$$\Delta_\mu^{en}(\mu) := \left(\langle \mathcal{L}_m(\mu)e^0, e^0 \rangle + \frac{\Delta t}{\alpha_{LB}(\mu)} \sum_{i=0}^{K-1} \|\mathcal{R}^i(\mu)\|^2 \right)^{1/2},$$

where $\alpha_{LB}(\mu)$ is a computable lower bound of the coercivity constant of $\mathcal{L}_a$, i.e. $0 < \alpha_{LB}(\mu) \leq \alpha(\mu)$.

Proof. The proof [24] makes repeated use of Young's inequality, i.e. for all $\varepsilon, a, b \in \mathbb{R}$ holds $ab \leq \frac{1}{2\varepsilon}a^2 + \frac{1}{2}\varepsilon^2b^2$. Starting with the error evolution equation (55), making use of the additive decomposition (56) and taking the scalar product with e^{k+1} yields

$$\langle \mathcal{L}_m e^{k+1}, e^{k+1} \rangle + \Delta t \langle \mathcal{L}_a e^{k+1}, e^{k+1} \rangle = \langle \mathcal{L}_m e^k, e^{k+1} \rangle + \Delta t \langle \mathcal{R}^k, e^{k+1} \rangle. \quad (59)$$

The first term on the right hand side can be bounded by Cauchy-Schwartz and Young's inequality with $\varepsilon = 1$:

$$\begin{aligned} \langle \mathcal{L}_m e^k, e^{k+1} \rangle &\leq \langle \mathcal{L}_m e^k, e^k \rangle^{1/2} \langle \mathcal{L}_m e^{k+1}, e^{k+1} \rangle^{1/2} \\ &\leq \frac{1}{2} \langle \mathcal{L}_m e^k, e^k \rangle + \frac{1}{2} \langle \mathcal{L}_m e^{k+1}, e^{k+1} \rangle. \end{aligned}$$

The second term on the right hand side of (59) can be bounded by Young's inequality with $\varepsilon^2 = \alpha$ and coercivity:

$$\begin{aligned} \Delta t \langle \mathcal{R}^k, e^{k+1} \rangle &\leq \Delta t \|\mathcal{R}^k\| \|e^{k+1}\| \\ &\leq \Delta t \left(\frac{1}{2\alpha} \|\mathcal{R}^k\|^2 + \frac{1}{2} \alpha \|e^{k+1}\|^2 \right) \\ &\leq \Delta t \left(\frac{1}{2\alpha} \|\mathcal{R}^k\|^2 + \frac{1}{2} \langle \mathcal{L}_a e^{k+1}, e^{k+1} \rangle \right). \end{aligned}$$

Then, (59) implies

$$\frac{1}{2} \langle \mathcal{L}_m e^{k+1}, e^{k+1} \rangle - \frac{1}{2} \langle \mathcal{L}_m e^k, e^k \rangle + \frac{1}{2} \Delta t \langle \mathcal{L}_a e^{k+1}, e^{k+1} \rangle \leq \Delta t \frac{1}{2\alpha(\mu)} \|\mathcal{R}^k\|^2.$$

Summation over $k = 0, \dots, K$ yields a telescope sum and simplifies to

$$\frac{1}{2} \langle \mathcal{L}_m e^K, e^K \rangle - \frac{1}{2} \langle \mathcal{L}_m e^0, e^0 \rangle + \frac{1}{2} \Delta t \sum_{k=1}^K \langle \mathcal{L}_a e^k, e^k \rangle \leq \sum_{k=0}^{K-1} \frac{\Delta t}{2\alpha} \|\mathcal{R}^k\|^2.$$

Multiplication with 2 and adding the e^0 -term gives the statement. $\square$

Remark 3.12 (Extensions). *Extensions of the error estimator $\Delta^{en}(\mu)$ exist. For example, $\mathcal{L}_l$ may be non-coercive [48], or $\mathcal{L}_E$ may be Δt -dependent [31], i.e. one can allow also explicit discretization contributions as obtained in Euler forward or Crank Nicolson time discretization. More error estimators can be derived by a space-time (Petrov-) Galerkin viewpoint, cf. [67, 74].*

Remark 3.13 (Output Estimation). *We did not yet address output estimation for the instationary case. This can be realized in simple or in more advanced ways, similar to the approaches of Sec. 2. Possible outputs in time-dependent scenarios can be time-dependent outputs $s^k(\mu)$ at each time step or a single scalar quantity $s(\mu)$ for the complete time-trajectory. For this, the problem $(E(\mu))$ can be extended by $l^k \in X', k = 0, \dots, K$ and*

$$s^k(\mu) = l^k(u^k; \mu), \quad k = 0, \dots, K, \quad s(\mu) = \sum_{k=0}^K s^k(\mu).$$

Then, one possibility for output estimation is the direct extension of the procedure of Def. 2.9: The reduced problem can be extended by

$$s_N^k(\mu) := l^k(u_N^k; \mu), \quad k = 0, \dots, K, \quad s_N(\mu) := \sum_{k=0}^K s_N^k(\mu).$$

Then, output error bounds are obtained using continuity of the output functionals:

$$\begin{aligned} |s^k(\mu) - s_N^k(\mu)| &\leq \Delta_s^k(\mu) := \|l^k(\cdot; \mu)\|_{X'} \Delta_u^k(\mu), \quad k = 0, \dots, K, \\ |s(\mu) - s_N(\mu)| &\leq \Delta_s(\mu) := \sum_{k=0}^K \Delta_s^k(\mu). \end{aligned}$$

This procedure again is admittedly very coarse. First, as only a “linear” dependence on the state error bound is obtained, and secondly, as the possible bad effectivity of the state bounds is inherited to the output bounds. Using a primal-dual technique, better estimates of single outputs can be obtained, cf. [24]. The dual problem to a scalar output of an instationary problem is a backward-in-time problem, where the inhomogeneities are given by the output functionals. Then, similar to Def. 2.26, an output correction can be performed by the primal residual applied to the dual solution, and output error bounds can be obtained which have the “squared” effect as in (20).

3.4 Offline/Online Decomposition

The offline/online decomposition is analogous to the stationary case. The main insight is that with time-independent operator components, the offline storage does not grow with K , but is independent of the time-step number. Herewith we again obtain an online phase, which is independent of the dimension H of the spatial discretization.

We again assume that $X = \text{span}(\psi_i)_{i=1}^{\mathcal{N}}$ is a discrete high-dimensional space of dimension $\mathcal{N}$, we are given the inner product matrix $\mathbf{K}$ and system matrix and vector components

$$\begin{aligned} \mathbf{K} &:= (\langle \psi_i, \psi_j \rangle)_{i,j=1}^{\mathcal{N}} \in \mathbb{R}^{\mathcal{N} \times \mathcal{N}}, & \mathbf{b}_q &:= (\langle b_q, \psi_i \rangle)_{i=1}^{\mathcal{N}} \in \mathbb{R}^{\mathcal{N}}, \\ \mathbf{L}_{I,q} &:= (\langle \mathcal{L}_{I,q} \psi_j, \psi_i \rangle)_{i,j=1}^{\mathcal{N}} \in \mathbb{R}^{\mathcal{N} \times \mathcal{N}}, & \mathbf{L}_{E,q} &:= (\langle \mathcal{L}_{E,q} \psi_j, \psi_i \rangle)_{i,j=1}^{\mathcal{N}} \in \mathbb{R}^{\mathcal{N} \times \mathcal{N}} \end{aligned} \quad (60)$$

for $q = 1, \dots, Q_b$, and Q_I, Q_E , respectively, and $\mathbf{u}_q^0 = (u_{q,i}^0)_{i=1}^{\mathcal{N}} \in \mathbb{R}^{\mathcal{N}}$, $q = 1, \dots, Q_u^0$ the coefficient vector of $u_q^0 = \sum_{i=1}^{\mathcal{N}} u_{q,i}^0 \psi_i \in X$. For a solution of $(E(\mu))$, one can assemble the system components by evaluating the coefficient functions and linear combinations

$$\begin{aligned} \mathbf{u}^0(\mu) &= \sum_{q=1}^{Q_u^0} \theta_{u^0,q}(\mu) \mathbf{u}_q^0, & \mathbf{b}^k(\mu) &= \sum_{q=1}^{Q_b} \theta_{b,q}^k(\mu) \mathbf{b}_q, \\ \mathbf{L}_I^k(\mu) &= \sum_{q=1}^{Q_I} \theta_{I,q}^k(\mu) \mathbf{L}_{I,q}, & \mathbf{L}_E^k(\mu) &= \sum_{q=1}^{Q_E} \theta_{E,q}^k(\mu) \mathbf{L}_{E,q}, \end{aligned}$$

and iteratively solve

$$\mathbf{L}_I^k(\mu) \mathbf{u}^{k+1} = \mathbf{L}_E^k(\mu) \mathbf{u}^k + \mathbf{b}^k(\mu), \quad k = 0, \dots, K-1, \quad (61)$$

in order to obtain the vector of unknowns $\mathbf{u}^k = (u_i^k)_{i=1}^{\mathcal{N}} \in \mathbb{R}^{\mathcal{N}}$ of $u^k = \sum_{i=1}^{\mathcal{N}} u_i^k \psi_i \in X$.

Remark 3.14 (Time Evolution for Non-Variational Discretizations). *For variational discretizations, e.g. finite elements, the matrices $\mathbf{L}_{I,q}, \mathbf{L}_{E,q}$ are exactly the components of the FEM mass or stiffness matrices, hence assembly is clear. For non-variational discretizations such as Finite Differences or Finite Volumes one can apply the current technique by assuming a quadrature approximation of the $L^2(\Omega)$ scalar product based on the current grid points, cf. Example 8 and 9. In both cases, $\mathbf{K}$ will be a diagonal matrix. The discretization for FD or FV schemes is frequently not given in terms of the above variational matrices, but matrices and vectors of unknowns are given by point-evaluations of the operator results, i.e. if $v = \mathcal{L} u$ for $u, v \in X$, then a linear operator $\mathcal{L} : X \rightarrow X$ is realized by a matrix operation $\mathbf{v} = \hat{\mathbf{L}} \mathbf{u}$. The relation to a matrix $\mathbf{L} = (\langle \mathcal{L} \psi_j, \psi_i \rangle)_{i,j=1}^{\mathcal{N}}$ of type (60) is simple: Set $v = \mathcal{L} \psi_j$, then $\mathbf{v} = \hat{\mathbf{L}} \mathbf{e}_j$ and*

$$(\mathbf{L})_{ij} = \mathbf{e}_i^T \mathbf{L} \mathbf{e}_j = \langle \mathcal{L} \psi_j, \psi_i \rangle = \langle v, \psi_i \rangle = \langle \psi_i, v \rangle = \mathbf{e}_i^T \mathbf{K} \mathbf{v} = \mathbf{e}_i^T \mathbf{K} \hat{\mathbf{L}} \mathbf{e}_j.$$

Hence $\mathbf{L} = \mathbf{K}\hat{\mathbf{L}}$. Therefore, the evolution step (61) reads

$$\mathbf{K}\hat{\mathbf{L}}_I^k(\mu)\mathbf{u}^{k+1} = \mathbf{K}\hat{\mathbf{L}}_E^k(\mu)\mathbf{u}^k + \mathbf{K}\hat{\mathbf{b}}^k(\mu), \quad k = 0, \dots, K-1,$$

hence, one can also omit $\mathbf{K}$ in the evolution for the solution of $(E(\mu))$. In particular this procedure is implemented in RBmatlab and used in the experiments of this section.

Now, the offline-online decomposition of $(E_N(\mu))$ is straightforward:

Proposition 3.15 (Offline/Online Decomposition of $(E_N(\mu))$).

(Offline Phase:) After computation of a reduced basis $\Phi_N = \{\varphi_1, \dots, \varphi_N\}$ compute the parameter- and time-independent matrices and vectors

$$\begin{aligned} \mathbf{b}_{N,q} &:= (\langle b_q, \varphi_i \rangle)_{i=1}^N \in \mathbb{R}^N, \quad \mathbf{L}_{N,I,q} := (\langle \mathcal{L}_{I,q} \varphi_j, \varphi_i \rangle)_{i,j=1}^N \in \mathbb{R}^{N \times N}, \\ \mathbf{u}_{N,q}^0 &:= (\langle u_q^0, \varphi_i \rangle)_{i=1}^N \in \mathbb{R}^N, \quad \mathbf{L}_{N,E,q} := (\langle \mathcal{L}_{E,q} \varphi_j, \varphi_i \rangle)_{i,j=1}^N \in \mathbb{R}^{N \times N}. \end{aligned}$$

(Online Phase:) For a given $\mu \in \mathcal{P}$ evaluate the coefficient functions $\theta_{I,q}^k(\mu)$, $\theta_{E,q}^k(\mu)$, $\theta_{u^0}(\mu)$ and $\theta_b^k(\mu)$, assemble the reduced system matrices and vectors

$$\begin{aligned} \mathbf{L}_{N,I}^k(\mu) &:= \sum_{q=1}^{Q_I} \theta_{I,q}^k(\mu) \mathbf{L}_{N,I,q}, \quad \mathbf{L}_{N,E}^k(\mu) := \sum_{q=1}^{Q_E} \theta_{E,q}^k(\mu) \mathbf{L}_{N,E,q}, \\ \mathbf{b}_N^k(\mu) &:= \sum_{q=1}^{Q_b} \theta_b^k(\mu) \mathbf{b}_{N,q}, \quad k = 0, \dots, K-1 \end{aligned}$$

and solve the discrete reduced evolution system by $\mathbf{u}_N^0 := \sum_{q=1}^{Q_{u^0}} \theta_{u^0,q}(\mu) \mathbf{u}_{N,q}^0$ and

$$\mathbf{L}_{N,I}^k(\mu) \mathbf{u}_N^{k+1} = \mathbf{L}_{N,E}^k(\mu) \mathbf{u}_N^k + \mathbf{b}_N^k(\mu), \quad k = 0, \dots, K-1.$$

Again, the computational procedure for obtaining the components is very simple: If we again assume the reduced basis to be given as coefficient matrix $\Phi_N \in \mathbb{R}^{N \times N}$, the components can be computed by

$$\mathbf{L}_{N,E,q} := \Phi^T \mathbf{L}_{E,q} \Phi, \quad \mathbf{L}_{N,I,q} := \Phi^T \mathbf{L}_{I,q} \Phi, \quad \mathbf{b}_{N,q} := \Phi^T \mathbf{b}_q, \quad \mathbf{u}_{N,q}^0 := \Phi^T \mathbf{u}_q^0.$$

The offline/online decomposition of the error estimators can also be realized. Computational procedures for obtaining upper continuity and lower coercivity bounds have been addressed in Sec. 2.5. The remaining ingredient for $\Delta_u^k(\mu)$ or Δ_u^{en} is again an efficient computational procedure for the residual norm. Analogous to Prop. 2.30 we obtain parameter-separability of the residual.

Proposition 3.16 (Parameter-Separability of the Residual). Set $Q_R := N(Q_E + Q_I) + Q_b$ and define $\mathcal{R}_q \in X$, $q = 1, \dots, Q_R$ by

$$\begin{aligned} (\mathcal{R}_1, \dots, \mathcal{R}_{Q_R}) &:= (\mathcal{L}_{E,1} \varphi_1, \dots, \mathcal{L}_{E,Q_E} \varphi_1, \dots, \mathcal{L}_{E,1} \varphi_N, \dots, \mathcal{L}_{E,Q_E} \varphi_N, \\ &\quad \mathcal{L}_{I,1} \varphi_1, \dots, \mathcal{L}_{I,Q_I} \varphi_1, \dots, \mathcal{L}_{I,1} \varphi_N, \dots, \mathcal{L}_{I,Q_I} \varphi_N, b_1, \dots, b_{Q_b}). \end{aligned}$$

Let $u_N^k(\mu) = \sum_{i=1}^N u_{N,i}^k \varphi_i$, $k = 0, \dots, K$ be the solution of $(E_N(\mu))$ and define $\theta_{\mathcal{R}}^k(\mu) := (\theta_{\mathcal{R},q}^k(\mu))_{k=0}^{K-1}$ by

$$\begin{aligned} (\theta_{\mathcal{R},1}^k(\mu), \dots, \theta_{\mathcal{R},Q_R}^k(\mu)) := & \frac{1}{\Delta t} \left(\theta_{E,1}^k(\mu) u_{N,1}^k(\mu), \dots, \theta_{E,Q_E}^k(\mu) u_{N,1}^k(\mu), \dots, \right. \\ & \theta_{E,1}^k(\mu) u_{N,N}^k(\mu), \dots, \theta_{E,Q_E}^k(\mu) u_{N,N}^k(\mu), \\ & -\theta_{I,1}^k(\mu) u_{N,1}^{k+1}(\mu), \dots, -\theta_{I,Q_I}^k(\mu) u_{N,1}^{k+1}(\mu), \dots, \\ & \left. -\theta_{I,1}^k(\mu) u_{N,N}^{k+1}(\mu), \dots, -\theta_{I,Q_I}^k(\mu) u_{N,N}^{k+1}(\mu), \right. \\ & \left. \theta_{b,1}^k(\mu), \dots, \theta_{b,Q_b}^k(\mu) \right). \end{aligned}$$

Then the residual $\mathcal{R}^k(\mu)$ defined in (54) is parameter-separable with

$$\mathcal{R}^k(\mu) = \sum_{q=1}^{Q_R} \theta_{\mathcal{R},q}^k(\mu) \mathcal{R}_q.$$

Herewith the norm $\|\mathcal{R}^k(\mu)\|$ can be computed efficiently with the offline/online procedure as stated in Prop. 2.32, now using the Gramian matrix $\mathbf{G}_R := (\langle \mathcal{R}_q, \mathcal{R}_{q'} \rangle)_{q,q'=1}^{Q_R}$, the coefficient vector $\theta_{\mathcal{R}}^k(\mu)$ and computing

$$\|\mathcal{R}^k(\mu)\| = \sqrt{\theta_{\mathcal{R}}^k(\mu)^T \mathbf{G}_R \theta_{\mathcal{R}}^k(\mu)}.$$

This completes the offline/online computational procedure for the general RB-approach for instationary problems.

We close this section on offline-online decomposition with a remark concerning possible coupling of time-step and spatial discretization.

Remark 3.17 (Coupling of Spatial and Time Discretization). *Note, that a slight dependence of the online phase on the spatial discretization exists via possible time-step constraints: The reduced problem has identical number of timesteps as the full problem. If Δt is constrained to $\mathcal{O}(\Delta x)$ for explicit discretizations of advection terms or $\mathcal{O}(\Delta x^2)$ for explicit discretization of diffusion operators, this implies that the full problem cannot be chosen highly accurate without affecting the time-discretization and herewith the reduced simulation. So in RB-approaches of time-evolution problems, only the complexity due to spatial, but not the time-discretization is reduced.*

3.5 Basis Generation

Again, we address the basis generation as separate section, although it is naturally part of the offline phase.

The most simple basis type for the instationary case is obtained by considering time as an additional “parameter” and then using a Lagrangian reduced basis according to Def. 2.39, i.e. $\Phi_N = \{u^{k(i)}(\mu^{(i)})\}_{i=1}^N$.

While this procedure is good for validation purposes, cf. Prop. 3.7, or for testing an RB-scheme, there are serious difficulties with this approach for obtaining a *good and small* basis. First, the time-parameter manifold may be more complex and hence many snapshots and herewith a larger basis may be required. Further, it is unclear, how to choose the time-indices $k^{(i)}$, and technically, for obtaining a single $u^{k^{(i)}}(\mu^{(i)})$ one needs to compute the complete trajectory $u^k(\mu^{(i)})$ for $k = 0, \dots, k^{(i)}$, and discard the unused information of the initial time steps.

The first procedure we will present addresses the first difficulty, the treatment of large snapshot sizes. The so called *Proper Orthogonal Decomposition (POD)* allows to compress large snapshot sets to the most important *POD modes*, that means a few vectors or functions containing the most important information of the data. Starting with a large number of functions $\{u_i\}_{i=1}^n \subset X$ the POD generates a small orthonormal set of basis functions Φ_N with $N \ll n$ by means of the so called empirical correlation operator. Technically, the POD corresponds to the Principal Component Analysis [42], the Karhunen-Loeve [45, 50] or the Hotelling Transformation [38]. We restrict ourselves to the definition and some elementary properties and refer to [72, 42] for details.

Proposition 3.18 (Proper Orthogonal Decomposition). *Let $\{u_i\}_{i=1}^n \subset X$ be a given set of snapshots. Then define the empirical correlation operator $R : X \rightarrow X$ by*

$$Ru := \frac{1}{n} \sum_{i=1}^n \langle u_i, u \rangle u_i, \quad u \in X.$$

Then R is a compact self-adjoint linear operator and there exists an orthonormal set $\{\varphi_i\}_{i=1}^{n'}$ of $n' \leq n$ eigenvectors with real eigenvalues $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_{n'} > 0$ with

$$Ru = \sum_{i=1}^{n'} \lambda_i \langle \varphi_i, u \rangle \varphi_i. \quad (62)$$

We denote $\Phi_N := \text{POD}_N(\{u_i\}) := \{\varphi_i\}_{i=1}^N$ as the POD-basis of size $N \leq n'$.

Proof. R is linear and bounded by $\|R\| \leq \frac{1}{n} \sum_{i=1}^n \|u_i\|^2$ and has finite-dimensional range, so R is compact. Further, R is self adjoint, as

$$\langle Ru, v \rangle = \frac{1}{n} \sum_i \langle u_i, u \rangle \langle u_i, v \rangle = \langle u, Rv \rangle,$$

hence, by the spectral theorem there exists an orthonormal system satisfying the spectral decomposition (62). It must be finite, $n' < \infty$, as the range of R is finite. $\square$

Some interesting properties are the following, see also Fig. 11 for an illustration.

- $\{\varphi_i\}$ as orthonormal basis is not unique due to sign change in each vector, or possible rotations if an eigenspace has dimension larger than one.
- The bases obtained by POD are hierarchical, i.e $\Phi_{N'} \subseteq \Phi_N$ for $N' \leq N$.

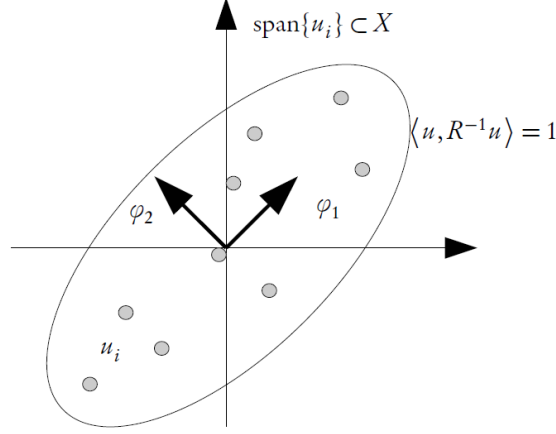


Figure 11: Illustration of POD.

- The POD does not depend on the order of the data $\{u_i\}_{i=1}^n$ (in contrast to a Gram-Schmidt orthonormalization).
- Let $X_{\text{POD},N} := \text{span}(\text{POD}_N(\{u_i\}_{i=1}^n))$. Then φ_1 is the direction of highest variance of $\{u_i\}_{i=1}^n$, φ_2 is the direction of highest variance of the projected data $\{P_{X_{\text{POD},1}}^\perp u_i\}_{i=1}^n$, etc.
- The coordinates of the data with respect to the POD-basis are uncorrelated, cf. Exercise 5.18.
- $\{\varphi_i\}$ and $\{\sqrt{\lambda_i}\}$ are the principal axis and axis intercepts of the ellipsoid $\langle u, R^{-1}u \rangle = 1$.
- The POD has a best-approximation property with respect to the squared error and the error can be exactly computed by the truncated eigenvalues, cf. Exercise 5.19.

$$\inf_{\substack{Y \subset X \\ \dim(Y)=N}} \frac{1}{n} \sum_{i=1}^n \|u_i - P_Y u_i\|^2 = \frac{1}{n} \sum_{i=1}^n \|u_i - P_{X_{\text{POD},N}} u_i\|^2 = \sum_{i=N+1}^{n'} \lambda_i. \quad (63)$$

The eigenvalue problem for the correlation operator either is very high-dimensional ($\dim X = \mathcal{N}$) or even infinite-dimensional ($\dim(X) = \infty$). This poses challenges for computational procedures. If the number of snapshots $n < \mathcal{N}$, the $\mathcal{N}$ -dimensional eigenvalue problem can be reformulated by an n -dimensional eigenvalue problem for the Gramian matrix and hence provide a more efficient computational procedure. Such a reformulation is sometimes referred to as *method of snapshots* [66] or the *kernel trick* in machine learning [64]. We leave the proof of the following proposition as Exercise 5.20.

Proposition 3.19 (Computation by Gram matrix). *Let $\mathbf{K}_u := (\langle u_i, u_j \rangle)_{i,j=1}^n \in \mathbb{R}^{n \times n}$ be the Gramian matrix of the snapshot set $\{u_i\}_{i=1}^n \subset X$. Then the following are equivalent*

- i) $\varphi \in X$ is an eigenvector of R for eigenvalue $\lambda > 0$ with norm 1 and a representation $\varphi = \sum_i a_i u_i$ with $\mathbf{a} = (a_i)_{i=1}^n \in \ker(\mathbf{K}_u)^\perp$
- ii) $\mathbf{a} = (a_i)_{i=1}^n \in \mathbb{R}^n$ is an eigenvector of $\frac{1}{n} \mathbf{K}_u$ with eigenvalue $\lambda > 0$ and norm $\frac{1}{\sqrt{n\lambda}}$.

Remark 3.20 (Difference between Greedy and POD). *We briefly comment on some differences between the POD of Prop. 3.18 and the strong greedy procedure of Def. 2.40, i.e. using the true projection error as indicator, $\Delta(Y, \mu) := \|u(\mu) - P_Y u(\mu)\|$. Both methods require the set of snapshots to be available, hence many full simulations. The main difference is the error measure that is guiding both procedures. The POD is aiming at minimizing the mean squared projection error, while the greedy procedure is aiming at minimizing the maximum projection error. So, “outliers” with single large error are allowed in POD, while the greedy algorithm will prevent such large deviations. Computationally, the greedy procedure produces an RB-space spanned by snapshots, i.e. a Lagrangian RB-space, the POD produces a space that is subset of a span of snapshots, but it is not a Lagrangian RB-space.*

Now, the greedy and POD procedure can be suitably combined to produce incrementally good bases for time-dependent problems. The resulting algorithm is called POD-Greedy procedure (initially denoted PCA-fixspace in [31]) and meanwhile is standard in RB-approaches for time-dependent problems, cf. [17, 19]. The idea is to “be greedy” with respect to the parameter and use POD with respect to time: We search the currently worst resolved parameter using an error bound or indicator $\Delta(Y, \mu)$, then compute the complete trajectory of the corresponding solution, orthogonalize this trajectory to the current RB-space, perform a POD with respect to time in order to compress the error trajectory to its most important new information, and add the new POD-mode to the current basis.

The use of the POD in the POD-Greedy procedure eliminates the two remaining problems that were stated above: We do not need and worry how to select time-instants for basis extension, and we do not discard valuable information in the computed trajectory, but try to extract maximal new information from the selected trajectory.

Definition 3.21 (POD-Greedy Procedure). *Let $S_{\text{train}} \subset \mathcal{P}$ be a given training set of parameters and $\varepsilon_{\text{tol}} > 0$ a given error tolerance. Set $X_0 := \{0\}$, $\Phi_0 := \emptyset$, $n := 0$ and define*

iteratively

$$\begin{aligned}
\text{while } \varepsilon_n &:= \max_{\mu \in S_{\text{train}}} \Delta(X_n, \mu) > \varepsilon_{\text{tol}} \\
\mu^{(n+1)} &:= \arg \max_{\mu \in S_{\text{train}}} \Delta(X_n, \mu) \\
u_{n+1}^k &:= u^k(\mu^{(n+1)}), k = 0, \dots, K, \text{ solution of } (E(\mu^{(n+1)})) \\
e_{n+1}^k &:= u_{n+1}^k - P_{X_n} u_{n+1}^k, k = 0, \dots, K \\
\varphi_{n+1} &:= \text{POD}_1(\{e_{n+1}^k\}_{k=0}^K) \\
\Phi_{n+1} &:= \Phi_n \cup \{\varphi_{n+1}\} \\
X_{n+1} &:= X_n + \text{span}(\varphi_{n+1}) \\
n &\leftarrow n + 1 \\
\text{end while.}
\end{aligned} \tag{64}$$

Thus, the output of the algorithm is the desired RB-space X_N and basis Φ_N by setting $N := n + 1$ as soon as (64) is false. The algorithm can be used with different error measures, e.g. the true squared projection error $\Delta(Y, \mu) := \sum_{k=0}^K \|u^k(\mu) - P_Y u^k(\mu)\|^2$ or one of the error estimators $\Delta(Y, \mu) := \Delta_u^K(\mu)$ or $\Delta(Y, \mu) := \Delta_u^{en}(\mu)$. In the first case, we denote the algorithm the *strong POD-Greedy* procedure, while in the latter cases it is called a *weak POD-Greedy* procedure.

Remark 3.22 (Consistency of POD-Greedy to POD and Greedy). *In the case of a single timestep $K = 1$, the evolution scheme corresponds to a single system solve and, hence, can be interpreted as solving a stationary problem. The POD-Greedy procedure then just is the standard greedy algorithm with included orthogonalization. On the other hand, it can easily be seen that for the case of a single parameter $\mathcal{P} = \{\mu\} \subset \mathbb{R}^p$ but arbitrary $K > 1$ the POD and POD-Greedy coincide: The POD-Greedy incrementally finds the next orthogonal basis vector, which is just the next POD-mode. Therefore, the POD-Greedy basis of a single trajectory is optimal in the least-squares sense. Thus, the POD-Greedy is a generalization of the POD.*

Again, the procedure is not only heuristic, but indeed also convergence rates can be derived, and a result analogous to Prop. 2.43 holds after some extensions [28]. First, a space-time norm for trajectories $v := (v^k(\mu))_{k=0}^K \subset X^{K+1}$ is introduced by suitable weights $w_k > 0, \sum_k w_k = T$ and

$$\|v\|_T := \left(\sum_{k=0}^K w_k \|v^k\|^2 \right)^{1/2}.$$

Then, the manifold of parametric trajectories is introduced as

$$\mathcal{M}_T := \{u(\mu) = (u^k(\mu))_{k=0}^K | \mu \in \mathcal{P}\} \subset X^{K+1},$$

while the *flat* manifold is

$$\mathcal{M} := \{u^k(\mu) | k = 0, \dots, K, \mu \in \mathcal{P}\} \subset X.$$

The componentwise projection on a subspace $Y \subset X$ is defined by $P_{T,Y} : X^{K+1} \rightarrow Y^{K+1}$ via

$$P_{T,Y} u := (P_Y u^k)_{k=0}^K, \quad u \in X^{K+1}.$$

Now the convergence rate statement can be formulated, where again d_n denotes the Kolmogorov n -width as defined in (37).

Proposition 3.23 (Convergence Rates of POD-Greedy Procedure). *Let $S_{\text{train}} = \mathcal{P}$ be compact, and the error indicator Δ chosen such that for suitable $\gamma \in (0, 1]$ holds*

$$\left\| u(\mu^{(n+1)}) - P_{T,X_n} u(\mu^{(n+1)}) \right\|_T \geq \gamma \sup_{u \in \mathcal{M}_T} \left\| u - P_{T,X_n} u \right\|_T. \quad (65)$$

i) (Algebraic convergence rate:) if $d_n(\mathcal{M}) \leq M n^{-\alpha}$ for some $\alpha, M > 0$ and all $n \in \mathbb{N}$ and $d_0(\mathcal{M}) \leq M$ then

$$\varepsilon_n \leq C M n^{-\alpha}, n > 0$$

with suitable (explicitly computable) constant $C > 0$.

ii) (Exponential convergence rate:) if $d_n(\mathcal{M}) \leq M e^{-a n^2}$ for $n \geq 0$, $M, a, \alpha > 0$ then

$$\varepsilon_n \leq C M e^{-c n^\beta}, n \geq 0$$

with $\beta := \alpha/(\alpha + 1)$ and suitable (explicitly computable) constants $c, C > 0$.

Remark 3.24 (Choice of Initial Basis). *The POD-Greedy procedure can also be used in a slightly different fashion, namely as extension of an existing basis $\Phi \neq \emptyset$. For this we simply set $N_0 := |\Phi|$, the initial basis $\Phi_{N_0} := \Phi$ and RB-space $X_{N_0} := \text{span}(\Phi)$ and start the POD-Greedy with N_0 instead of 0. One possible scenario for this could be the improvement of an existing basis: If insufficient accuracy is detected in the online phase, one can return to the offline phase for a basis extension. Another useful application of this variant is choosing the initial basis such that $u_q^0 \in X_{N_0}$. Then the RB-error at time $t = 0$ is zero and the a-posteriori error bounds are more tight, as the initial error contribution is zero.*

Remark 3.25 (Adaptivity, PT-Partition). *As the solution complexity of time-dependent simulations is much higher than for stationary problems, the size of S_{train} is more critical in the instationary case: Although the a-posteriori error estimators have complexity polynomial in N they still are not very cheap due to the linear scaling with K . The size of S_{train} being limited, the choice and adaptation of the training set of parameters, [29, 37] becomes an even more important issue for instationary problems.*

Especially for time-dependent problems, the solution variability with varying parameter and time can be very dramatic (e.g. a transport process with varying directions, see the model problem of Sec. 3.1). Then the parameter-domain partitioning approaches mentioned in Sec. 2.6 can also be applied. In particular the h p-RB-approach has been extended to parabolic problems [19] and the P-partition approach has been used for hyperbolic problems [29].

If the time-interval is large such that solution variation are too high, parameter domain partitioning may be not sufficient. In these cases, also time-interval partitioning can be applied [14, 16]. This means that single bases are constructed for subintervals of the time-axis. For the reduced simulation, a suitable switching between the spaces must be realized over time. This can either be done during the Galerkin-projection step of the RB-scheme, or done as a separate orthogonal projection step. In both cases, the error estimation procedure can be kept fully rigorous by suitably incorporating the additional projection errors [14]. The division of the time-interval can be obtained adaptively: Starting with a large interval the basis generation with the POD-Greedy is initiated. As soon as a too large basis is obtained (or anticipated early by extrapolation), the POD-Greedy is terminated, the time interval is split and separate bases are constructed on the subintervals.

We also want to conclude this section with some experiments which again can be reproduced with the script `rb_tutorial.m`. We choose the model problem of Sect. 3.1 and apply a finite volume discretization. For this we assume a uniform hexaedral grid consisting of 64×32 squares, and set $\Delta t = 1/256$. The advection is discretized explicitly with an upwind flux and the diffusion is discretized implicitly, cf. [31]. The time-step width is sufficiently small in order to meet the CFL timestep restriction. The snapshot plots in Fig. 10 already were based on this discretization.

We generate a POD-Greedy basis using the initial data field u^0 as starting basis, setting $\varepsilon_{\text{tol}} = 1 \cdot 10^{-2}$ and choosing the X -norm error indicator from Prop. 3.9 as selection criterion $\Delta(Y, \mu) := \Delta_{\mu}^K(\mu)$ with the (coarse) bound constants $\gamma_{\text{UB}}(\mu) = 1$ and $\alpha_{\text{LB}}(\mu) = 1$, as the explicit and inverted implicit spatial discretization operators are L^2 -stable due to our choice of time-step width. As training parameter set we use the vertices of a uniform 10×10 grid of points on $\mathcal{P} = [0, 1]^2$.

The resulting POD-Greedy training estimator development is illustrated in Fig. 12a) by plotting for each basis size N the maximal estimator over the training set of parameters. It nicely shows an exponentially decaying behavior, however the convergence is much slower compared to the stationary case due to the complex parameter and time-dependence. In fact the low-diffusion region is very hard to approximate, which causes the relatively large basis sizes. This difficult region is also reflected in Fig. 12b) where parameter selection frequencies are plotted over the training set. Larger circles indicate training parameters which are chosen more frequently during basis generation. Note, that this is a difference to the greedy algorithm in the stationary case: Parameters can be selected multiple times during the POD-Greedy procedure, as addition of a single mode to the basis does not necessarily reduce the error to zero.

The first 16 generated basis vectors are plotted in Fig. 13. One can observe increasing oscillations in the basis functions. One can also observe that the initial condition, cf. Fig. 10, was chosen as initial basis.

In Fig. 14a) we illustrate the behavior of the a-posteriori error bound Δ_{μ}^k and the error

$$e^k(\mu) := \max_{k'=0, \dots, k} \left\| u^{k'}(\mu) - u_N^{k'}(\mu) \right\|$$

at final time $k = K$ for a parameter sweep along the diagonal of $\mathcal{P}$ by $\mu = s \cdot$

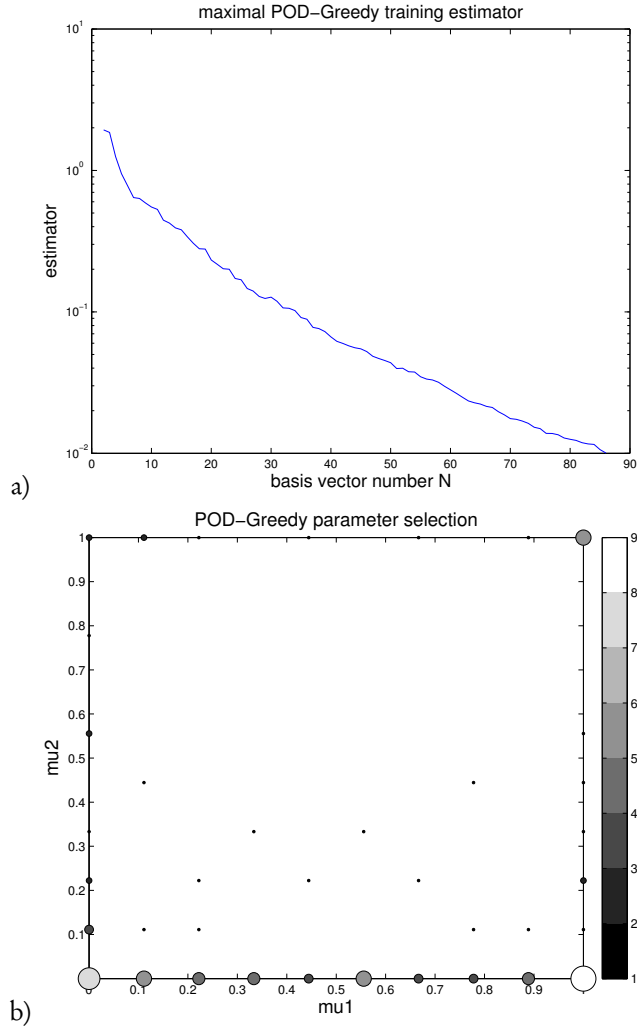


Figure 12: Illustration of POD-Greedy results for advection-diffusion model problem. a) Plot of maximal training estimator decay and b) plot of parameter selection frequency.

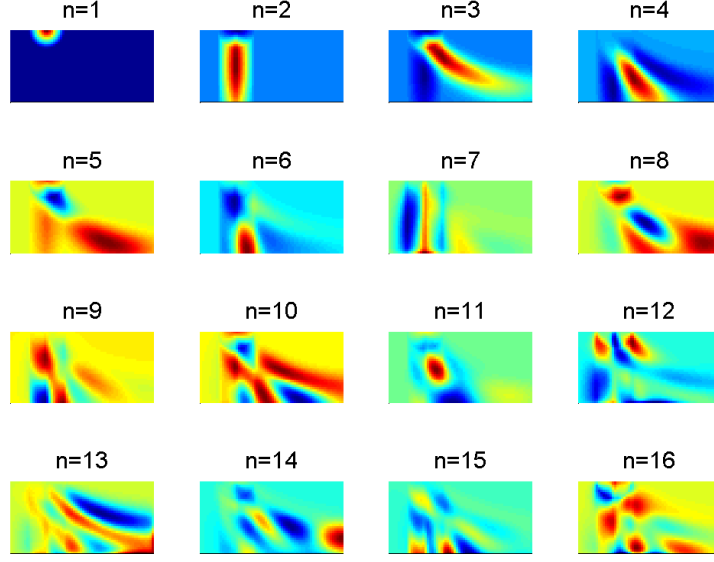


Figure 13: Illustration of the first 16 basis vectors produced by the POD-Greedy procedure.

$(1, 1)^T, s \in [0, 1]$. It can be verified that indeed the error estimator is below ε_{tol} for the training points obtained by $s = i/10, i = 0, \dots, 10$. This bound cannot be guaranteed for test-points in the low-diffusivity region. The results indicate that it would be beneficial to include more training points in this difficult parameter region. For example, one could choose the diffusivity values log-equidistant in accordance with Prop. 2.45 or apply an adaptive training set extension algorithm [29, 30]. The ratio of the error bound and the true error is about one order of magnitude, hence can be considered to be quite good. This is made more explicit in Fig. 14b), where for a test set of 200 random parameters, we plot the effectivities $\eta^k(\mu) := \Delta_\mu^k(\mu)/e^k(\mu)$ at final time $k = K$ where the test-points are sorted according to μ_2 . We nicely see that the effectivities are lower bounded by 1, i.e. the error estimators are reliable, while the factor of overestimation is not too large. Note, however, that the error estimators are mostly incremental with growing k , hence the effectivities are expected to get worse for larger times T . This can be improved by space-time Galerkin approaches and estimators [67, 74].

4 Extensions and Outlook

We give some comments and references to literature on further aspects and some current developments that could not be covered by this introductory tutorial chapter.

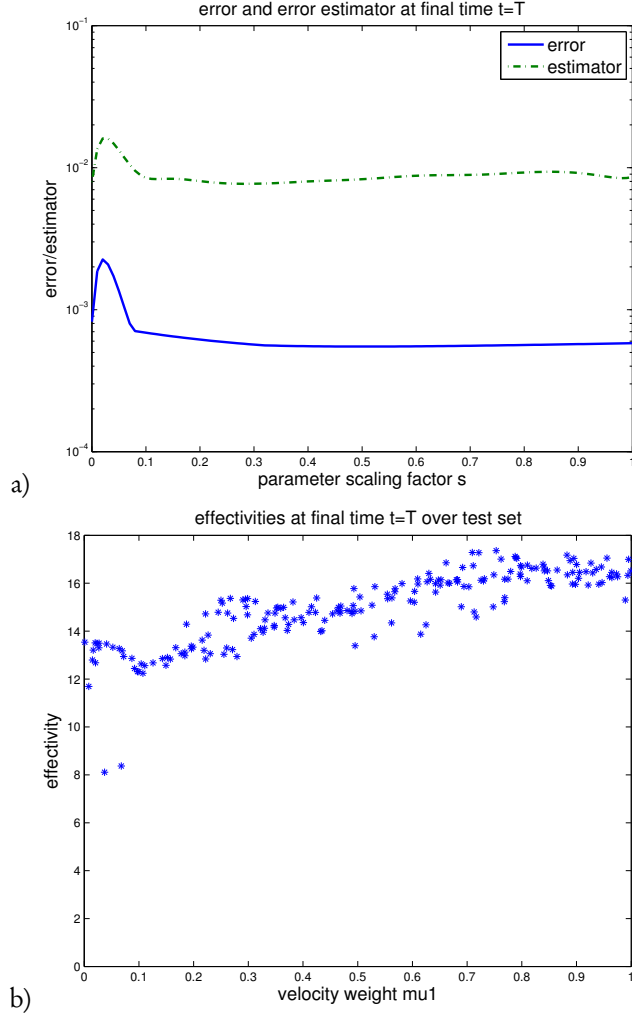


Figure 14: Behavior of error estimator $\Delta_{\mu}^k(\mu)$ and error $e^k(\mu)$ at end time $k = K$. a) Error and estimator over parameter sweep along the diagonal of the parameter domain, $\mu = s \cdot (1, 1)^T, s \in [0, 1]$, b) effectivities $\eta^K(\mu)$ for 200 random parameter vectors.

Some extensions are meanwhile well-established. The first question is the treatment of nonlinear problems. In the case of simple “polynomial” nonlinearities, which can be written as a multilinear form in the variational form of the PDE, this multilinearity can be effectively used for suitable offline/online decomposition of the Galerkin-reduced system [69] and the Newton-type iteration for solving the fixpoint equation. Also, a-posteriori error analysis is possible for these RB-approaches making use of the Brezzi-Rappaz-Raviart theory. Problems that can be treated by this are nonlinear diffusion or nonlinear advection problems, e.g. the Burgers equation [70]. For more general nonlinearities, the EI-method can be applied for approximating stationary and instationary nonlinear problems [25]. A specialization of this procedure for discretization operators is the Empirical Operator Interpolation initially used in [33] and then extended to nonlinear problems in [32, 17], which requires local reconstruction of the reduced solution and local evaluation of the differential operator. This procedure has later been denoted Discrete Empirical Interpolation Method (DEIM) [10] in the context of nonlinear state-space dynamical systems. Note, however, that the stability of the approximated RB-systems involving an EI approximation step is a nontrivial aspect.

The second obvious possibility for extensions of the presented methodology is the treatment of more general linear problems. As harmonic Maxwell’s or Helmholtz equations are non-coercive, a more general notion of stability is considered in RB-approaches, the *inf-sup* stability. This notion generalizes the coercivity (inf-sup constant being always at least as large as coercivity constant) while the RB-error bounds have frequently identical structure and the inf-sup constant “replaces” the coercivity constant [61, 71]. A main problem is that inf-sup stability is not inherited to subspaces, hence separate test- and trial spaces must be constructed in order to guarantee stability of the reduced systems. By suitable definition of a norm on the test space, optimal stability factors can be obtained by double greedy procedures [11]. In the context of time-dependent problems, an interesting possibility is the formulation as a time-space variational form [67], which represents an inf-sup stable formulation. The resulting RB error estimators are typically very sharp in contrast to the incremental estimators of Sec. 3. In particular they have provable effectivity bounds. However, the discretization then must be adjusted to be able to cope with the additional dimension by the time variable. The notion of inf-sup stability is also required in treatment of systems of PDEs, most notably the Stokes system [62] for viscous flow. Parameter dependence can also be obtained by geometry parametrization [63]. As a general solution strategy, the parametrized PDE is mapped to a reference domain which incorporates the geometry parameters in the coefficient functions of the transformed PDE.

Various recent developments can be found, which represent active research directions in RB-methods. First, the variational problems can be additionally constrained with inequalities. The resulting variational inequalities can also be treated successfully with RB-methods, both in stationary and instationary cases [23, 34, 35]. The Stokes system reveals a saddle point structure which is typical for other types of problems. Such general saddle point problems have been considered in [22]. Extensions to more complex coupled systems, e.g. Navier-Stokes [13, 69] or the Boussinesq-approximation [47], can be found. An important field of application for RB-methods

are multi-query scenarios in optimization or optimal control. Both parameter optimization [15, 56] as well as optimal control for elliptic and parabolic equations have been investigated [12, 43, 44]. Parameters can also be considered as random influence in PDEs, corresponding Monte-Carlo approaches using reduced order models can be realized [7, 36]. Multiscale problems which require multiple evaluation of micro-models for a macro-scale simulation, can make use of RB-micromodels [6, 48], or domain partitioning approaches can be used to capture global parametric information in local bases [46]. An important branch of past and current application are domain decomposition approaches based on RB-models for simple geometries, which then are coupled to more complex geometrical shapes involving a huge number of local parameters. The original Reduced Basis Element method [53] has been extended to an extremely flexible static condensation approach [39] which allows to construct online various geometries by using RB-models as “Lego” building blocks. Iterative domain decomposition schemes [54] can be used for handling distributed coupled problems.

5 Exercises

Exercise 5.1 (Finite-Dimensional X_N for Thermal Block). *Show that the thermal block model for $B_1 = 1$ has a solution manifold, which is contained in an $N := B_2$ -dimensional linear subspace $X_N \subset H_{\Gamma_D}^1$. This can be obtained by deriving an explicit solution representation. In particular, find N snapshot parameters $\mu^{(i)}, i = 1, \dots, N$, such that $X_N = \text{span}\{u(\mu^{(1)}), \dots, u(\mu^{(N)})\}$.*

Exercise 5.2 (Many Parameters, Simple Solution). *Devise an instantiation of $(P(\mu))$ with arbitrary number of parameters $p \in \mathbb{N}$ but solution being contained in a 1-dimensional linear subspace.*

Exercise 5.3 (Conditions for Uniform Coercivity). *Assume a parameter-separable bilinear form $a(\cdot, \cdot; \mu)$. Under which conditions on the coefficient functions θ_q^a and the components a_q can uniform coercivity of a be concluded?*

Exercise 5.4 (Thermal Block as Instantiation of $(P(\mu))$). *Verify that the thermal block satisfies the assumption of $(P(\mu))$, i.e. uniform continuity, coercivity and parameter-separability of the bilinear and linear forms. In particular, specify the constants, the coefficient functions and components. Argue that it even is an example for a compliant problem, i.e. symmetric and $f = l$.*

Exercise 5.5 (Finite-Dimensional Exact Approximation for $Q_a = 1$). *Assume a general problem of type $(P(\mu))$ with parameter-separable forms. Show that in the case of $Q_a = 1$ the solution manifold $\mathcal{M}$ is contained in a reduced basis space X_N of dimension at most Q_f , and show that there exist $N \leq Q_f$ parameters $\mu^{(i)}, i = 1, \dots, N$ such that*

$$X_N = \text{span}(u(\mu^{(1)}), \dots, u(\mu^{(N)})).$$

Exercise 5.6 (Lipschitz-Continuity of $(P(\mu))$). *Let $\theta_q^a, \theta_q^f, \theta_q^l$ be Lipschitz-continuous with respect to μ with Lipschitz-constants L_q^a, L_q^f, L_q^l .*

i) Show that a, f, l are Lipschitz continuous by computing suitable constants L_a, L_f, L_l such that

$$\begin{aligned} |a(u, v; \mu) - a(u, v; \mu')| &\leq L_a \|u\| \|v\| \|\mu - \mu'\|, \\ |f(v; \mu) - f(v; \mu')| &\leq L_f \|v\| \|\mu - \mu'\|, \\ |l(v; \mu) - l(v; \mu')| &\leq L_l \|v\| \|\mu - \mu'\|, \quad u, v \in X, \mu, \mu' \in \mathcal{P}. \end{aligned}$$

ii) Derive suitable constants L_u, L_s such that

$$\|u(\mu) - u(\mu')\| \leq L_u \|\mu - \mu'\|, \quad |s(\mu) - s(\mu')| \leq L_s \|\mu - \mu'\|, \quad \mu, \mu' \in \mathcal{P}.$$

$$\text{Hint: } L_u = L_f / \bar{\alpha} + \bar{\gamma}_f L_a / \bar{\alpha}^2, \quad L_s = L_l \bar{\gamma}_f / \bar{\alpha} + \bar{\gamma}_l L_u.$$

Remark: An identical statement holds for the reduced solutions $u_N(\mu), s_N(\mu)$.

Exercise 5.7 (Differentiability of $u(\mu)$). Prove Prop. 2.8.

Exercise 5.8 (Best-approximation Bound for Symmetric Case). Show that if $a(\cdot, \cdot; \mu)$ is symmetric, for all $\mu \in \mathcal{P}$ holds

$$\|u(\mu) - u_N(\mu)\| \leq \sqrt{\frac{\gamma(\mu)}{\alpha(\mu)}} \inf_{v \in X_N} \|u(\mu) - v\|.$$

(This is a sharpening of (4), and hence the Lemma of Céa by a square root.)

Exercise 5.9 (Relative Error and Effectivity Bounds). Provide proofs for the relative a-posteriori error estimate and effectivity, Prop. 2.22.

Exercise 5.10 (Energy Norm Error and Effectivity Bounds). Provide proofs for the energy norm a-posteriori error estimate and effectivity, Prop. 2.23.

Exercise 5.11 ($\alpha(\mu)$ for Thermal Block). Show for the thermal block that $\alpha(\mu) = \min_i \mu_i$. (Therefore, $\alpha_{\text{LB}}(\mu)$ can be chosen as that value in the error bounds.) Similarly, show that $\gamma(\mu) = \max_i \mu_i$. (Hence effectivities are always bounded by $\mu_{\max} / \mu_{\min}$ if $\mu \in [\mu_{\min}, \mu_{\max}]^p$.)

Exercise 5.12 (Max-theta Approach for Continuity Upper Bound). Let a be symmetric and all a_q positive semidefinite and $\theta_q^a(\mu) > 0, q = 1, \dots, Q_a, \mu \in \mathcal{P}$. Let $\bar{\mu} \in \mathcal{P}$ and $\gamma(\bar{\mu})$ be known. Show that for all $\mu \in \mathcal{P}$ holds

$$\gamma(\mu) \leq \gamma_{\text{UB}}(\mu) < \infty,$$

with continuity upper bound

$$\gamma_{\text{UB}}(\mu) := \gamma(\bar{\mu}) \max_{q=1, \dots, Q_a} \frac{\theta_q^a(\mu)}{\theta_q^a(\bar{\mu})}.$$

Exercise 5.13 (Monotonicity of Greedy Error). *Prove that the greedy algorithm produces monotonically decreasing error sequences $(\varepsilon_n)_{n \geq 1}$ if*

- i) $\Delta(Y, \mu) := \|u(\mu) - P_Y u(\mu)\|$, i.e. the orthogonal projection error is chosen as error indicator.
- ii) we have the compliant case ($a(\cdot, \cdot; \mu)$ symmetric and $f(\cdot; \mu) = l(\cdot; \mu)$) and $\Delta(Y, \mu) := \Delta_u^{en}(\mu)$, i.e. the energy error estimator from Prop. 2.23 is chosen as error indicator.

Exercise 5.14 (Gramian Matrix and Properties). *Let $\{u_1, \dots, u_n\} \subset X$ be a finite subset. Define the Gramian matrix through*

$$\mathbf{G} := (\langle u_i, u_j \rangle)_{i,j=1}^n \in \mathbb{R}^{n \times n}.$$

Show that the following holds:

- i) $\mathbf{G}$ is symmetric and positive semidefinite,
- ii) $\text{rank}(\mathbf{G}) = \dim(\text{span}(\{u_i\}_{i=1}^n))$,
- iii) $\{u_i\}_{i=1}^n$ are linearly independent $\Leftrightarrow \mathbf{G}$ is positive definite.

(Recall that such Gramian matrices already appeared as $\mathbf{G}_l, \mathbf{G}_r, \mathbf{K}_N$ in the offline/online decomposition in Prop. 2.32, Prop. 2.33 and Prop. 2.34.)

Exercise 5.15 (Gram Schmidt Orthonormalization). *Prove that the procedure given in Prop. 2.48 indeed produces the Gram-Schmidt orthonormalized sequence.*

Exercise 5.16 (Orthonormalization of Reduced Basis). *Prove that the procedure given in Prop. 2.48 produces an orthonormal basis also if $\mathbf{C}$ is chosen differently, as long as it satisfies $\mathbf{C}\mathbf{C}^T = \mathbf{G}^{-1}$. (Hence, more than only Cholesky-factorization is possible.)*

Exercise 5.17 (Uniform Boundedness with Respect to Δt). *Prove the statement of Prop. 3.4. Hint:*

$$\left(\frac{1}{1 + \alpha \frac{T}{K}} \right)^K = \left(\left(\frac{1}{1 + \frac{\alpha T}{K}} \right)^{\frac{K}{\alpha T}} \right)^{\alpha T} \rightarrow e^{-\alpha T}$$

as $K \rightarrow \infty$.

Exercise 5.18 (Data Uncorrelated in POD-Coordinates). *Verify that the coordinates of the data $\{u_i\}_{i=1}^n$ with respect to the POD-basis $\{\varphi_j\}_{j=1}^{n'}$ are uncorrelated and the mean squared coordinates are just the eigenvalues of the correlation operator, i.e. for all $j, k = 1, \dots, n', j \neq k$ holds*

$$\sum_{i=1}^n \langle u_i, \varphi_j \rangle \langle u_i, \varphi_k \rangle = 0, \quad \frac{1}{n} \sum_{i=1}^n \langle u_i, \varphi_j \rangle^2 = \lambda_j.$$

Exercise 5.19 (POD Mean Squared Error). *Prove the second equality in (63): Let $\{u_i\}_{i=1}^n$ be given. Show that the mean squared error of the POD-projection can be explicitly obtained by the sum of the truncated eigenvalues, i.e.*

$$\frac{1}{n} \sum_{i=1}^n \|u_i - P_{X_{\text{POD},N}} u_i\|^2 = \sum_{i=N+1}^{n'} \lambda_i.$$

Exercise 5.20 (POD via Gramian Matrix). *Prove Prop. 3.19, i.e. the equivalence of computation of the POD via the correlation operator or the Gramian matrix.*

Acknowledgements

The author wants to acknowledge the Baden-Württemberg Stiftung gGmbH for funding as well as the German Research Foundation (DFG) for financial support within the Cluster of Excellence in Simulation Technology (EXC 310/1) at the University of Stuttgart.

References

- [1] B. ALMROTH, P. STERN, AND F. BROGAN, *Automatic choice of global shape functions in structural analysis*, AIAA J., 16 (1978), pp. 525–528. (Cited on p. 46)
- [2] M. BARRAULT, Y. MADAY, N. NGUYEN, AND A. PATERA, *An ‘empirical interpolation’ method: application to efficient reduced-basis discretization of partial differential equations*, C. R. Math. Acad. Sci. Paris Series I, 339 (2004), pp. 667–672. (Cited on p. 40)
- [3] R. BECKER AND R. RANNACHER, *Weighted a-posteriori error estimates in FE methods*, in Proc. ENUMATH-97, 1998, pp. 621–637. (Cited on p. 20)
- [4] G. BERKOOZ, P. HOLMES, AND J. L. LUMLEY, *The proper orthogonal decomposition in the analysis of turbulent flows*, Annual Review of Fluid Mechanics, 25 (1993), pp. 539–575. (Cited on p. 46)
- [5] P. BINEV, A. COHEN, P. DAHMEN, R. DEVORE, G. PETROVA, AND P. WJASZCZYK, *Convergence rates for greedy algorithms in reduced basis methods*, SIAM J. Math. Anal., 43 (2011), pp. 1457–1472. (Cited on p. 34)
- [6] S. BOYVAL, *Reduced-basis approach for homogenization beyond the periodic setting*, SIAM Multiscale Modeling and Simulation, 7 (2008), pp. 466–494. (Cited on p. 69)
- [7] S. BOYVAL, C. L. BRIS, Y. MADAY, N. NGUYEN, AND A. PATERA, *A reduced basis approach for variational problems with stochastic parameters: Application to heat conduction with variable Robin coefficient*, Computer Methods in Applied Mechanics and Engineering, 1 (2009). (Cited on p. 69)

- [8] D. BRAESS, *Finite Elemente - Theorie, schnelle Löser und Anwendungen in der Elastizitätstheorie*, 4. Aufl., Springer, 2007. (Cited on p. 8)
- [9] A. BUFFA, Y. MADAY, A. PATERA, C. PRUD'HOMME, AND G. TURINICI, *A priori convergence of the greedy algorithm for the parametrized reduced basis*, ESAIM M2AN, Math. Model. Numer. Anal., 46 (2012), pp. 595–603. (Cited on p. 34)
- [10] S. CHATURANTABUT AND D. SORESENSEN, *Nonlinear Model Reduction via Discrete Empirical Interpolation*, SIAM J. Sci. Comput., 32 (2010), pp. 2737–2764. (Cited on p. 68)
- [11] W. DAHMEN, C. PLESKEN, AND G. WELPER, *Double greedy algorithms: reduced basis methods for transport dominated problems*, ESAIM Math. Model. Numer. Anal., 48 (2014), pp. 623–663. (Cited on p. 68)
- [12] L. DEDE, *Adaptive and Reduced Basis Methods for Optimal Control Problems in Environmental Applications*, PhD thesis, Doctoral School in Mathematical Engineering, Politecnico di Milano, 2008. (Cited on p. 69)
- [13] S. DEPARIS, *Reduced basis error bound computation of parameter-dependent Navier-Stokes equations by the natural norm approach*, SIAM J. Num. Anal., 46 (2008), pp. 2039–2067. (Cited on p. 68)
- [14] M. DIHLMANN, M. DROHMANN, AND B. HAASDONK, *Model reduction of parametrized evolution problems using the reduced basis method with adaptive time-partitioning*, in Proc. of ADMOS 2011, 2011. (Cited on p. 64)
- [15] M. A. DIHLMANN AND B. HAASDONK, *Certified PDE-constrained parameter optimization using reduced basis surrogate models for evolution problems*, COAP, Computational Optimization and Applications, (2014). (Cited on pp. 50, 69)
- [16] M. DROHMANN, B. HAASDONK, AND M. OHLBERGER, *Adaptive reduced basis methods for nonlinear convection-diffusion equations*, in Proc. FVCA6, 2011. (Cited on p. 64)
- [17] ———, *Reduced basis approximation for nonlinear parametrized evolution equations based on empirical operator interpolation*, SIAM J. Sci. Comput., 34 (2012), pp. A937–A969. (Cited on pp. 44, 61, 68)
- [18] J. EFTANG, M. GREPL, AND A. PATERA, *A posteriori error bounds for the empirical interpolation method*, C. R. Acad. Sci. Paris, Ser. I 348, (2010), pp. 575–579. (Cited on p. 43)
- [19] J. EFTANG, D. KNEZEVIC, AND A. PATERA, *An hp certified reduced basis method for parametrized parabolic partial differential equations*, Math. Comp. Model Dyn., 17 (2011), pp. 395–422. (Cited on pp. 36, 47, 61, 63)
- [20] J. L. EFTANG, A. T. PATERA, AND E. M. RØNQVIST, *An hp certified reduced basis method for parametrized elliptic partial differential equations*, SIAM J. Sci Comp, 32 (2010), pp. 3170–3200. (Cited on pp. 10, 36, 47)

- [21] J. FINK AND W. RHEINBOLDT, *On the error behaviour of the reduced basis technique for nonlinear finite element approximations*, ZAMM, 63 (1983), pp. 21–28. (Cited on pp. 3, 31)
- [22] A.-L. GERNER AND K. VEROY, *Certified reduced basis methods for parametrized saddle point problems*, SIAM J. Sci. Comput., (2013). accepted. (Cited on p. 68)
- [23] S. GLAS AND K. URBAN, *On noncoercive variational inequalities*, SIAM J. Numer. Anal., 52 (2014), pp. 2250–2271. (Cited on p. 68)
- [24] M. GREPL, *Reduced-basis Approximations and a Posteriori Error Estimation for Parabolic Partial Differential Equations*, PhD thesis, Massachusetts Institute of Technology, May 2005. (Cited on pp. 43, 46, 47, 53, 54, 55)
- [25] M. GREPL, Y. MADAY, N. NGUYEN, AND A. PATERA, *Efficient reduced-basis treatment of nonaffine and nonlinear partial differential equations*, ESAIM M2AN, Math. Model. Numer. Anal., 41 (2007), pp. 575–605. (Cited on pp. 47, 68)
- [26] M. GREPL AND A. PATERA, *A posteriori error bounds for reduced-basis approximations of parametrized parabolic partial differential equations*, M2AN, Math. Model. Numer. Anal., 39 (2005), pp. 157–181. (Cited on pp. 46, 47, 53)
- [27] B. HAASDONK, *Reduzierte-Basis-Methoden*, Lecture Notes, IANS-Report 4/11, University of Stuttgart, Germany, 2011. (Cited on pp. 4, 16, 17, 41)
- [28] B. HAASDONK, *Convergence rates of the POD-Greedy method*, ESAIM M2AN, Math. Model. Numer. Anal., 47 (2013), pp. 859–873. (Cited on p. 62)
- [29] B. HAASDONK, M. DIHLMANN, AND M. OHLBERGER, *A training set and multiple basis generation approach for parametrized model reduction based on adaptive grids in parameter space*, Math. Comp. Model. Dyn., 17 (2012), pp. 423–442. (Cited on pp. 36, 37, 63, 66)
- [30] B. HAASDONK AND M. OHLBERGER, *Basis construction for reduced basis methods by adaptive parameter grids*, in Proc. International Conference on Adaptive Modeling and Simulation, ADMOS 2007, P. Díez and K. Runesson, eds., CIMNE, Barcelona, 2007. (Cited on pp. 36, 66)
- [31] ———, *Reduced basis method for finite volume approximations of parametrized linear evolution equations*, ESAIM M2AN, Math. Model. Numer. Anal., 42 (2008), pp. 277–302. (Cited on pp. 46, 51, 55, 61, 64)
- [32] B. HAASDONK AND M. OHLBERGER, *Reduced basis method for explicit finite volume approximations of nonlinear conservation laws*, in Proc. HYP 2008, International Conference on Hyperbolic Problems: Theory, Numerics and Applications, vol. 67 of Proc. Sympos. Appl. Math., Providence, RI, 2009, Amer. Math. Soc., pp. 605–614. (Cited on p. 68)

- [33] B. HAASDONK, M. OHLBERGER, AND G. ROZZA, *A reduced basis method for evolution schemes with parameter-dependent explicit operators*, Electron. Trans. Numer. Anal., 32 (2008), pp. 145–161. (Cited on p. 68)
- [34] B. HAASDONK, J. SALOMON, AND B. WOHLMUTH, *A reduced basis method for parametrized variational inequalities*, SIAM J. Num. Anal., 50 (2012), pp. 2656–2676. (Cited on p. 68)
- [35] B. HAASDONK, J. SALOMON, AND B. WOHLMUTH, *A reduced basis method for the simulation of american options*, in Proc. ENUMATH 2011, 2012. (Cited on p. 68)
- [36] B. HAASDONK, K. URBAN, AND B. WIELAND, *Reduced basis methods for parametrized partial differential equations with stochastic influences using the Karhunen Loeve expansion*, SIAM/ASA J. Unc. Quant., 1 (2013), pp. 79–105. (Cited on p. 69)
- [37] J. S. HESTHAVEN, B. STAMM, AND S. ZHANG, *Efficient greedy algorithms for high-dimensional parameter spaces with applications to empirical interpolation and reduced basis methods*, ESAIM: M2AN, 48 (2014), pp. 259–283. (Cited on pp. 36, 63)
- [38] H. HOTELLING, *Analysis of a complex of statistical variables into principal components.*, J. Educ. Psychol., 24 (1933), pp. 417–441. (Cited on p. 59)
- [39] D. HUYNH, D. KNEZEVIC, AND A. PATERA, *A static condensation reduced basis element method: Approximation and a posteriori error estimation*, Mathematical Modelling and Numerical Analysis, 47 (2013), pp. 213–251. (Cited on p. 69)
- [40] D. HUYNH, G. ROZZA, S. SEN, AND A. PATERA, *A successive constraint linear optimization method for lower bounds of parametric coercivity and inf-sup stability constants*, C. R. Math. Acad. Sci. Paris Series I, 345 (2007), pp. 473–478. (Cited on pp. 29, 30)
- [41] K. ITO AND S. RAVINDRAN, *A reduced order method for simulation and control of fluid flows*, J. Comput. Phys., 143 (1998), pp. 403–425. (Cited on p. 46)
- [42] I. JOLLIFFE, *Principal Component Analysis*, Springer-Verlag, 2002. (Cited on p. 59)
- [43] M. KÄRCHER AND M. GREPL, *A certified reduced basis method for parametrized elliptic optimal control problems*, ESAIM: Control Optim. Calc. Var., 20 (2014), pp. 416–441. (Cited on p. 69)
- [44] ———, *A posteriori error estimation for reduced order solutions of parametrized parabolic optimal control problems*, ESAIM M2AN, Math. Model. Numer. Anal., 48 (2014), pp. 1615–1638. (Cited on pp. 47, 69)

- [45] K. KARHUNEN, *Über lineare Methoden in der Wahrscheinlichkeitsrechnung*, Ann. Acad. Sci. Fennicae, ser. A I. Math. Phys., 37 (1946). (Cited on p. 59)
- [46] S. KAULMANN, M. OHLBERGER, AND B. HAASDONK, *A new local reduced basis discontinuous Galerkin approach for heterogeneous multiscale problems*, C. R. Math. Acad. Sci. Paris, 349 (2011), pp. 1233–1238. (Cited on p. 69)
- [47] D. KNEZEVIC, N. NGUYEN, AND A. PATERA, *Reduced basis approximation and a posteriori error estimation for the parametrized unsteady Boussinesq equations*, Math. Mod. Meth. Appl. Sci., 21 (2011), pp. 1415–1442. (Cited on pp. 47, 68)
- [48] D. KNEZEVIC AND A. PATERA, *A certified reduced basis method for the Fokker-Planck equation of dilute polymeric fluids: FENE dumbbells in extensional flow.*, SIAM J. Sci. Comp., 32 (2010), pp. 793–817. (Cited on pp. 55, 69)
- [49] D. KRÖNER, *Numerical Schemes for Conservation Laws*, John Wiley & Sons and Teubner, 1997. (Cited on p. 50)
- [50] M. M. LOEVE, *Probability Theory*, Van Nostrand, Princeton, NJ, 1955. (Cited on p. 59)
- [51] Y. MADAY, N. NGUYEN, A. PATERA, AND G. PAU, *A general, multi-purpose interpolation procedure: the magic points*, Tech. Rep. RO7037, Laboratoire Jacques-Louis-Lions, Université Pierre et Marie Curie, Paris, 2007. (Cited on pp. 40, 41, 42, 43)
- [52] Y. MADAY, A. PATERA, AND G. TURINICI, *a priori convergence theory for reduced-basis approximations of single-parameter symmetric coercive elliptic partial differential equations*, C.R. Acad. Sci. Paris, Ser. I, 335 (2002), pp. 289–294. (Cited on p. 35)
- [53] Y. MADAY AND E. RÖNQVIST, *The reduced basis element method: Application to a thermal fin problem.*, SIAM J. Sci. Comput., 26 (2006), pp. 240–258. (Cited on p. 69)
- [54] I. MAIER AND B. HAASDONK, *A Dirichlet-Neumann reduced basis method for homogeneous domain decomposition*, Appl. Num. Math., 78 (2014), pp. 31–48. (Cited on p. 69)
- [55] A. NOOR AND J. PETERS, *Reduced basis technique for nonlinear analysis of structures*, AIAA J., 18 (1980), pp. 455–462. (Cited on p. 3)
- [56] I. OLIVEIRA AND A. PATERA, *Reduced-basis techniques for rapid reliable optimization of systems described by affinely parametrized coercive elliptic partial differential equations*, Optim. Eng., 8 (2007), pp. 43–65. (Cited on p. 69)
- [57] A. PATERA AND G. ROZZA, *Reduced Basis Approximation and a Posteriori Error Estimation for Parametrized Partial Differential Equations*, MIT, 2007. Version 1.0, Copyright MIT 2006-2007, to appear in (tentative rubric) MIT Papalardo Graduate Monographs in Mechanical Engineering. (Cited on pp. 4, 5, 6, 13, 16, 17, 27, 28, 35)

- [58] T. PORSCHING AND M. LEE, *The reduced basis method for initial value problems*, SIAM J. Numer. Anal., 24 (1987), pp. 1277–1287. (Cited on pp. 3, 46)
- [59] C. PRUD’HOMME, D. ROVAS, K. VEROY, L. MACHIELS, Y. MADAY, A. PATERA, AND G. TURINICI, *Reliable real-time solution of parametrized partial differential equations: Reduced-basis output bound methods*, J. Fluids Engineering, 124 (2002), pp. 70–80. (Cited on pp. 10, 20)
- [60] W. RHEINBOLDT, *On the theory and error estimation of the reduced-basis method for multi-parameter problems*, Nonlinear Analysis, Theory, Methods & Applications, 21 (1993), pp. 849–858. (Cited on p. 3)
- [61] D. ROVAS, *Reduced-Basis Output Bound Methods for Parametrized Partial Differential Equations*, PhD Thesis, MIT, Cambridge, MA, 2003. (Cited on p. 68)
- [62] G. ROZZA, *Shape design by optimal flow control and reduced basis techniques: Applications to bypass configurations in haemodynamics*, PhD Thesis, École Polytechnique Fédérale de Lausanne, November 2005. (Cited on p. 68)
- [63] G. ROZZA, D. HUYNH, AND A. PATERA, *Reduced Basis approximation and a posteriori error estimation for affinely parametrized elliptic coercive partial differential equations: application to transport and continuum mechanics*, Arch. Comput. Meth. Eng., 15 (2008), pp. 229–275. (Cited on pp. 4, 30, 68)
- [64] B. SCHÖLKOPF AND A. J. SMOLA, *Learning with Kernels: Support Vector Machines, Regularization, Optimization and Beyond*, MIT Press, 2002. (Cited on p. 60)
- [65] S. SEN, *Reduced basis approximations and a posteriori error estimation for many-parameter heat conduction problems*, Numerical Heat Transfer, Part B: Fundamentals, 54 (2008), pp. 369–389. (Cited on p. 36)
- [66] L. SIROVICH, *Turbulence and the dynamics of coherent structures. I. coherent structures*, Quart. Appl. Math., 45 (1987), pp. 561–571. (Cited on p. 60)
- [67] K. URBAN AND A. PATERA, *An improved error bound for reduced basis approximation of linear parabolic problems*, Math. Comput., 83 (2014), pp. 1599–1615. (Cited on pp. 47, 55, 66, 68)
- [68] K. URBAN, S. VOLKWEIN, AND O. ZEEB, *Greedy sampling using nonlinear optimization*, in *Reduced Order Methods for Modeling and Computational Reductions*, A. Quarteroni and G. Rozza, eds., Springer, 2014, pp. 139–157. (Cited on p. 36)
- [69] K. VEROY AND A. PATERA, *Certified real-time solution of the parametrized steady incompressible Navier-Stokes equations: Rigorous reduced-basis a posteriori error bounds*, Int. J. Numer. Meth. Fluids, 47 (2005), pp. 773–788. (Cited on p. 68)

- [70] K. VEROY, C. PRUD'HOMME, AND A. PATERA, *Reduced-basis approximation of the viscous Burgers equation: rigorous a posteriori error bounds*, C. R. Math. Acad. Sci. Paris Series I, 337 (2003), pp. 619–624. (Cited on p. 68)
- [71] K. VEROY, C. PRUD'HOMME, D. V. ROVAS, AND A. T. PATERA, *A posteriori error bounds for reduced-basis approximation of parametrized noncoercive and nonlinear elliptic partial differential equations*, in Proceedings of 16th AIAA computational fluid dynamics conference, 2003. Paper 2003-3847. (Cited on pp. 31, 68)
- [72] S. VOLKWEIN, *Model reduction using proper orthogonal decomposition*, lecture notes, University of Konstanz, 2013. (Cited on p. 59)
- [73] D. WIRTZ, D. SORENSEN, AND B. HAASDONK, *A-posteriori error estimation for DEIM reduced nonlinear dynamical systems*, SIAM J. Sci. Comput., 36 (2014), pp. A311–A338. (Cited on pp. 43, 44)
- [74] M. YANO AND A. PATERA, *A space-time variational approach to hydrodynamic stability theory*, Proceedings of the Royal Society A, 469 (2013). Article Number 20130036. (Cited on pp. 55, 66)

Index

- A-posteriori error bounds
 - energy norm, 17, 54, 70
 - output, 15, 21
 - relative, 16, 70
 - state space, 15, 21, 43, 46, 53
- Advection-diffusion model, 47
- Certification, *see* A-posteriori error bounds
- Coercivity constant, lower bound, 27, 29
- Coercivity, uniform, 7, 48
- Collateral basis, 41
- Compliant problem, 17
- Continuity constant, upper bound, 29, 70
- Continuity, uniform, 7, 48
- Convergence
 - El convergence rates, 43
 - global exponential convergence, 35
 - Greedy convergence rates, 34
 - POD-Greedy convergence rates, 63
 - uniform convergence, 14
- Differentiability, 10, 70
- Effectivities, 16, 17, 21
- Empirical correlation operator, 59
- Empirical Interpolation (EI)
 - a-posteriori error estimation, 43
 - a-priori convergence rate, 43
 - conservation property, 44
 - Lebesgue constant bound, 43
 - method, 40
- Error estimators, *see* A-posteriori error bounds
- Error indicator, choice of, 32
- Error-residual relation, 13, 52
- Evolution problem, 48
- Finite Difference discretization, 51
- Finite Volume discretization, 50
- Greedy
 - convergence rates, 34
 - multistage, 35
- procedure, 31
- strong, 34
- weak, 33, 34
- Kolmogorov width, 34
- Lagrangian Reduced Basis, 31
- Lipschitz continuity, 10, 69
- Magic points, 41
- Min-theta procedure, 28
- Multi-query, 1, 69
- Offline phase, *see* Offline/Online decomposition
- Offline/Online decomposition, 2, 22, 23, 26, 27, 56, 57
- Online phase, *see* Offline/Online decomposition
- Orthonormalization, 12, 37
- Output approximation, 11, 20, 55
- Overfitting, 33
- Parameter domain partitioning, 36
- Parameter separability, 8, 25, 49, 57
- Parameter vector, 2
- Parametric PDEs, 2
- POD-Greedy
 - convergence rates, 63
 - procedure, 61
- Primal-dual approach, 20
- Proper Orthogonal Decomposition (POD), 59
- RBmatlab, 4, 18, 37, 57
- Real-time, 1
- Reduced Basis
 - method, 2
 - space, 10
- Reproduction of solutions, 14, 52
- Slim computing, 2
- Snapshots, 2, 10
- Solution manifold, 2, 6, 8, 9

Space-time energy norm, 54
Stability, 8, 12, 49
State approximation, 11, 51
Successive Constraint Method (SCM), 29

Taylor Reduced Basis, 31
Thermal Block model, 5, 9, 69
Training set
 adaptivity, 36
 treatment, 35

Vanishing error bound, 15, 53