

Multigrid and Domain Decomposition Methods in Fault-Prone Environments

by
Christian Alexander Glusa

Ingénieur diplômé de l'École Polytechnique, 2011
Diplom-Mathematiker, Karlsruhe Institute of Technology, 2012
Sc. M., Brown University, 2013

A dissertation submitted in partial fulfillment of the
requirements for the Degree of Doctor of Philosophy
in the Division of Applied Mathematics at Brown University

Providence, Rhode Island

May 2017

© Copyright 2017 by Christian Alexander Glusa

This dissertation by Christian Alexander Glusa is accepted in its present form by
the Division of Applied Mathematics as satisfying the dissertation requirement
for the degree of Doctor of Philosophy.

Date _____

Mark Ainsworth, Director

Recommended to the Graduate Council

Date _____

George Em Karniadakis, Reader

Date _____

Johnny Guzmán, Reader

Approved by the Graduate Council

Date _____

Andrew G. Campbell
Dean of the Graduate School

In memoriam Vytas.

Vita

Education

2012 – 2017 PhD, Brown University, Providence, RI, USA

Advisor: Mark Ainsworth

2013 Masters of Science in Applied Mathematics, Brown University, Providence, RI, USA

2009 – 2011 Ingénieur polytechnicien, Ecole Polytechnique, Palaiseau, France

Double-degree program

Major: Applied Mathematics

Minors: Pure Mathematics, Physics, Mechanics, Computer Science

2006 – 2009 & 2011 – 2012, Diplom-Mathematiker, Karlsruhe Institute of Technology, Germany

Major: Mathematics

Minors: Theoretical Physics, Economics

Diploma thesis: Sedimentation in Periodic Domains

Advisor: Willy Dörfler

Talks & Presentations

- 3/2017, SIAM Conference on Computational Science and Engineering, Atlanta, *"The Resiliency of Multilevel Methods on Next Generation Computing Platforms: Probabilistic Model and Its Analysis"*
- 1/2017, CSRI, Sandia National Lab, Albuquerque, *"Resilience Analysis for Multigrid Methods"*
- 1/2017, Applied Math and Scientific Computing Seminar, Temple University, *"Resilience Analysis for Multigrid Methods"*
- 1/2017, Frontiers in Applied and Computational Mathematics, ICERM, Providence, *"Resilience Analysis for Multigrid Methods"* (Poster)
- 7/2016, SIAM Annual Meeting, Boston, *"Resilience Analysis for Multigrid Methods"* (Poster)
- 4/2016, CRUNCH seminar, Brown University, *"Resilience Analysis for Multigrid Methods"*
- 4/2016, Applied Math Days, Rensselaer Polytechnic Institute, *"Multigrid methods, random matrices and fault resilience"*
- 3/2016, 14th Copper Mountain Conference on Iterative Methods, *"Resilience Analysis for Multigrid Methods"*
- 11/2015, Math Slam, Brown University, *"Multigrid methods, random matrices and fault resilience"*
- 8/2015, 2015 AARMS Workshop on Domain Decomposition Methods for PDEs, Dalhousie University, Halifax, Canada, *"Multigrid methods, random matrices and fault resilience"*

- 7/2015, Math Slam, Brown University, *"Multigrid methods, random matrices and fault resilience"*
- 3/2015, Graduate student seminar, Brown University, *"Multigrid methods, random matrices and fault resilience"*
- 10/2011, Workshop for Particulate Flows, Karlsruhe Institute of Technology, Germany, *"Sedimentation in Periodic Domains"*

Other Conferences & Workshops Attended

- 8/2016, Argonne Training Program on Extreme-Scale Computing, St. Charles, Illinois
- 5/2016, HPC Day 2016, CSCVR, University of Massachusetts, Dartmouth
- 10/2015, Finite Element Circus, University of Massachusetts, Dartmouth
- 9/2015, Topical Workshop on Numerical Methods for Large-Scale Nonlinear Problems and Their Applications, ICERM, Providence
- 5/2014, Topical Workshop on Robust Discretization and Fast Solvers for Computable Multi-Physics Models, ICERM, Providence
- 8/2013, Brown-Kobe Summer School on High-Performance Computing, Kobe University, Japan

Teaching

- 2015–, Math Teacher at English for Action, Providence, RI, Weekly classes for immigrants on High School level

- Summer 2014, Scientific Computing Summer School, Brown University, Organized a three week summer school for undergrad students from Universit Pierre-et-Marie-Curie (Paris VI)
- Spring 2014, Teaching Assistant, Methods of Applied Mathematics II, Brown University
- Fall 2013, Teaching Assistant, Methods of Applied Mathematics I, Brown University

Service & Outreach

- Fall 2016, Judge for Brown Mathematical Contest for Modelling
- 2016,AWM Grad-Undergrad Mentor, Brown University
- Fall 2015, Judge for Brown Mathematical Contest for Modelling
- 2015–2016, Treasurer for the Brown University SIAM Student Chapter
- 2010–2011, Treasurer for Aslive, Versailles, France, Charitable society for adults with developmental disabilities
- Summer 2009, Summer camp organizer, Foyer Les Sources, Paris, France, Summer vacation camp for adults with developmental disabilities
- Summer 2008, Summer camp organizer, Foyer Les Sources, Paris, France, Summer vacation camp for adults with developmental disabilities
- 2005 – 2006, Voluntary social year, Foyer Les Sources, Paris, France, Care for adults with developmental disabilities

Honors

- 03/2014, Nominated to Sigma Xi Honor Society
- 08/2009–08/2011, Recipient of the "Jacques Garaialde Bernard Oppetit" scholarship, Excellence Scholarship
- 09/2010, Recipient of the "Prix du Projet Scientifique Collectif", Scientific group project of 9 months,
"Compréhension de la Nage à faible nombre de Reynolds".

Preface

Today, more than ever before, scientific exploration relies on large-scale computer simulation of complex phenomena. The installation of larger and faster supercomputers enables models of physical systems to include more features and make more accurate predictions. The advent of the new generation of exascale machines, as mandated by President Obama's executive order to establish the National Strategic Computing Initiative¹, poses several unique challenges. Historically, algorithms and their implementations were optimized with respect to the number of floating-point operations (flops) necessary. However, it has become apparent that the time to solution is more and more dependent on the amount of communication between the individual compute units of the system. A new constraint in the design of supercomputers is the power consumption of the installation. The energy necessary to send a message depends on distance of the communication, and simply up-scaling the design of current petascale systems, i.e. systems that can perform one petaflop, to the next generation of exascale machines is outright infeasible. The power consumption of such a system would easily surpass that of half a million homes. In order to observe the commonly accepted limit of 20MW, lower voltage thresholds and reduced cell capacitance are required. The inevitable consequence of these hardware level solutions are increased rates of faults that affect the correctness of the performed computations. This effect is amplified by the significant expansion

¹Creating a National Strategic Computing Initiative, Executive Order No. 13702 of July 29th, 2015, Federal Register

of the number of components making up an exascale system. Therefore, fault detection and mitigation strategies will have to be considered in order to take full advantage of the potential of the next generation of supercomputers [27, 28, 29, 42].

Many scientific simulations rely on the solution of very large, but sparse, linear systems of equations. Multigrid and Domain Decomposition methods are widely used solvers due to their attractive properties. Their benefits include optimal computational complexity, scaling and easy of implementation on parallel machines. It is therefore of prime concern to ensure the transition of these methods and their preexisting implementations to the next generation of high-performance systems and study their application to the problems that become tractable on these new machines.

The contribution of this work is two-fold. In the first part, we study the impact of faults on the convergence of iterative solvers of linear systems. The first chapter serves as an introduction to linear iterative solvers, and briefly describes one of the theoretical frameworks used to prove convergence of multigrid methods in the ideal, fault-free case. In Chapter 2, we describe the different categories of faults and how they are modeled using random matrices. We derive the error propagator of a fault-prone multigrid algorithm, which will be instrumental to prove convergence estimates in the presence of faults. Theoretical results concerning products of random matrices and methods of determining their asymptotic rate of convergence are given in Chapter 3. As it turns out, the so-called Lyapunov spectral radius is difficult to compute, but can be estimated in some cases. In Chapters 4 and 5, we apply the theoretical framework first to a simple iteration to illustrate the approach, and then to a two grid method. The extension to the multigrid case is given in Chapter 6. It will transpire that multigrid is not a fault-resilient solver, since convergence can only be achieved for limited problem size. By protecting one of its steps, the prolongation operation from coarser to finer meshes, fault resilience

is obtained. All theoretical results are illustrated by means of numerical examples that stem from a variety of different applications. More practical implementation details concerning the detection and mitigation of faults are discussed in Chapter 7. We illustrate how the theoretical results can be leveraged to obtain a fault-tolerant multigrid implementation. Finally, we discuss fault mitigation for overlapping domain decomposition solvers in Chapter 8. We extend the framework to account for faults that appear in unknowns that are shared between compute nodes, and give convergence estimates for one- and two-level methods.

In the second part of this work, we turn our attention to the efficient solution of non-local and fractional models using iterative linear solvers. Non-local equations, as compared to classical local models, can be used to more accurately describe the underlying physical phenomena. The price to pay is an increased complexity in terms of discretization and solution of the arising system of equations. We introduce archetypes of elliptic and parabolic non-local equations in Chapter 10, and state results concerning the regularity of their solution. We give details concerning finite element approximations and explain how optimal rates of convergence can be achieved through adaptive mesh refinement based on the use of a posteriori error estimators. In Chapter 11, we turn to the question of efficient implementations of assembly, solution and computation of error indicators. The finite element assembly involves the computation of singular element contributions for which special quadrature rules are required. Since the considered problems are non-local, the straightforward approach results in dense system matrices. Similarly, straightforward computation of a posteriori error estimators scales quadratically in the number of unknowns. We circumvent these issues by exploiting the low rank of off-diagonal blocks and explain how all steps of the assemble-solve-estimate-refine cycle can be achieved in quasi-linear complexity. Finally, in Chapter 12, we illustrate optimality of computational complexity and rate of convergence for a number of one- and two-dimensional numerical examples, involving fractional Poisson and

heat equations and a fractional reaction-diffusion system.

Acknowledgements

First and foremost, I want to thank my advisor Mark Ainsworth. This work would not have been possible without his guidance and support. I am very fortunate to have benefited from his insight and eye for interesting problems. I will miss our weekly meetings and his sense of humor.

I would also like to thank Professor George Em Karniadakis and Professor Johnny Guzmán for serving on my thesis committee and generously providing feedback about this work.

Finally, I want to acknowledge all the support I received from my colleagues, friends and family. You kept me sane throughout grad school, and I am grateful.

Contents

List of Tables	xix
List of Figures	xx
I Fault-Prone Iterative Methods	1
1 Multigrid Algorithm	2
1.1 Classical Iterative Solvers	3
1.2 Multigrid Algorithm	5
1.3 Convergence Theory for Multigrid	9
2 Fault-Prone Iterative Methods	12
2.1 Modeling Faults	14
2.2 Probabilistic Model of Faults	14
2.3 Scope of Faults Covered by the Analysis	17
2.4 A Model for Fault-Prone Multigrid	18
2.5 Faults in the Smoother	18
2.6 Faults in Restriction, Prolongation and Coarse Grid Correction	20
2.7 Model for Multigrid Algorithm in Presence of Faults	20
3 Lyapunov Exponents and Their Computation	24
3.1 Lyapunov Exponents	24

3.2	Monte-Carlo Approximation	32
3.3	Approximation via Invariant Distribution	32
3.4	Resampled Monte-Carlo Method	35
3.5	Replica Trick	36
4	Application of the Replica Trick to a Simple Iteration	39
5	Analysis of the Fault-Prone Two Grid Method	43
5.1	Numerical Verification of Non-Resilience	52
5.2	Protecting the Prolongation Restores Resilience	56
5.3	Protecting the Restriction Fails to Give Resilience	58
6	Analysis of the Fault-Prone Multigrid Method	62
6.1	The Effect of Protection of the Prolongation	69
6.2	Adaptively Refined Meshes	74
6.3	Higher Spatial Dimension and Higher Order Finite Elements	75
7	Implementation Issues	79
7.1	Detection of Soft Faults	79
7.2	Protection of the Prolongation	81
8	Fault-Prone Domain Decomposition Methods	86
8.1	Overlapping Domain Decomposition Methods	87
8.2	Modeling Blockwise Faults in Iterative Schemes	91
8.3	Convergence Analysis for Domain Decomposition Methods	94
8.4	Two-Level Methods	101
8.5	Discussion	111
9	Future Work on Fault Resilience and Related Topics	113
9.1	Fault-Resilient Domain Decomposition Methods	113
9.2	Asynchronous Iterations	114

II	Fast Solvers for Non-Local Equations	115
10	The Finite Element Method for the Fractional Laplacian	116
10.1	Fractional Laplacians	118
10.2	Weak formulation	120
10.3	Finite Element Approximation of the Fractional Poisson Equation . .	122
10.4	Adaptive Mesh Refinement	129
11	Overcoming Computational Challenges of the Fractional Laplacian	134
11.1	Computation of Entries of the Stiffness Matrix	134
11.1.1	Reduction to Smooth Integrals	135
11.1.2	Determining the Order of the Quadrature Rules	141
11.2	Solving the Linear Systems	146
11.3	Sparse Approximation of the Stiffness Matrix and of the Strong Form	150
12	Numerical Examples	158
12.1	Fractional Poisson Equation	158
12.1.1	One-Dimensional Examples	159
12.1.2	Two-Dimensional Examples	164
12.1.3	L-shaped Domain	170
12.1.4	Exploiting Symmetry	173
12.2	Fractional Heat Equation	175
12.3	Fractional Reaction-Diffusion Systems	179
12.4	Discussion	181
13	Future Work on Fractional and Non-Local Equations	184

★ Parts of this thesis have been previously published or have been submitted for publication:

- Mark Ainsworth and Christian Glusa. “Numerical Mathematics and Advanced

Applications - ENUMATH 2015: Proceedings of ENUMATH 2015”. In: Cham: Springer International Publishing, 2015. Chap. Multigrid at Scale?, pp. 237–253. ISBN: 978-3-319-39929-4. DOI: 10.1007/978-3-319-39929-4_24

- Mark Ainsworth and Christian Glusa. “Is the Multigrid Method Fault Tolerant? The Two-Grid Case”. In: *SIAM Journal on Scientific Computing* 39.2 (2017), pp. C116–C143. DOI: 10.1137/16M1100691
- Mark Ainsworth and Christian Glusa. “Is the multigrid method fault tolerant? The Multilevel Case.” In: *SIAM Journal on Scientific Computing* (Submitted)
- Mark Ainsworth and Christian Glusa. “Towards an Efficient Finite Element Method for the Integral Fractional Laplacian on Polygonal Domains”. In: *TBD* (2017)
- Mark Ainsworth and Christian Glusa. “Aspects of an adaptive finite element method for the fractional Laplacian: a priori and a posteriori error estimates, efficient implementation and multigrid solver”. Submitted

List of Tables

12.1 Complexity of different solvers for $(\boldsymbol{M} + \Delta t \boldsymbol{A}^s) \vec{u} = \vec{b}$ 176

12.2 IMEX scheme by Koto. 179

List of Figures

1.1	Multigrid hierarchy	8
2.1	Schematic representation of a node failure	19
5.1	Residual norm against iterations for Fault-Prone Two Grid Method .	53
5.2	Uniform mesh for the square domain.	54
5.3	Lyapunov spectral radius $\varrho(\mathcal{E}^{TG})$	55
5.4	Lyapunov spectral radius $\varrho(\mathcal{E}^{TG})$ with protected prolongation . . .	59
5.5	Residual against iterations for Fault-Prone Two Grid Method with protected prolongation	60
5.6	Lyapunov spectral radius $\varrho(\mathcal{E}^{TG})$ with protected restriction	61
6.1	Residual against iterations for the Fault-Prone W-Cycle Multigrid Method.	71
6.2	Lyapunov spectral radius $\varrho(\mathcal{E}_L)$ for the Fault-Prone W-Cycle Multi- grid method.	72
6.3	Residual against iterations for the Fault-Prone W-Cycle Multigrid Method with protected prolongation.	73
6.4	Initial mesh for the motor problem.	75
6.5	Lyapunov spectral radius $\varrho(\mathcal{E}_L)$ for the motor problem.	76
6.6	Mesh for the Fichera cube.	77
6.7	Lyapunov spectral radius $\varrho(\mathcal{E}_L)$ for the Fichera cube.	78

7.1	Lyapunov spectral radius $\varrho(\mathcal{E}_L)$ using $K = 2$ replicas in every operation.	81
7.2	Lyapunov spectral radius $\varrho(\mathcal{E}_L)$ using $K_P = 4$, $k_P = 3$ and $K = 3$. .	83
7.3	Level sets $\varrho(\mathcal{E}_L) = 0.57$ for different levels of detection and protection.	84
8.1	Asymptotic convergence rate $\varrho(\mathcal{E}_{ASH})$ for $N = 16$ subdomains. . .	100
8.2	Asymptotic convergence rate $\varrho(\mathcal{E}_{ASH})$	101
8.3	Asymptotic convergence rate $\varrho(\mathcal{E}_{ASH}\mathcal{E}_{CG})$ without protected prolongation for $N = 16$ subdomains.	102
8.4	Asymptotic convergence rate $\varrho(\mathcal{E}_{ASH}\mathcal{E}_{CG})$ without protected prolongation for $N \sim n$	103
8.5	Energy norm of the variance of the fault-prone coarse-grid correction.	109
8.6	Asymptotic convergence rate $\varrho(\mathcal{E}_{ASH}\mathcal{E}_{CG})$ with protected prolongation and $N = 16$ subdomains.	110
8.7	Asymptotic convergence rate $\varrho(\mathcal{E}_{ASH}\mathcal{E}_{CG})$ with protected prolongation and $N \sim n$ subdomains.	111
8.8	$\ N_{AS}\mathbf{A}\ _2$ in 1D for several mesh sizes h and $N = 4$ subdomains. .	112
10.1	Solutions u corresponding to the constant right-hand side for $s = 0.25$ and for $s = 0.75$	123
11.1	Element pairs containing a singularity.	135
11.2	Numbering of local nodes for touching triangular elements K and $\tilde{K}$ or element K and edge e	137
11.3	Cluster tree for a one-dimensional problem.	151
11.4	Cluster pairs for a one-dimensional problem.	152
11.5	Memory usage of the dense matrix and its sparse approximation. .	156
11.6	Assembly time of the dense matrix and its sparse approximation. . .	157

12.1	$\tilde{H}^s$ - and L^s -error for the one-dimensional problem with constant right-hand side.	160
12.2	Timings for assembly of the stiffness matrix, solution of linear system and computation of the error indicators for the one-dimensional problem with constant right-hand side.	161
12.3	Solutions to the one-dimensional fractional Poisson problem with right-hand side $f(x) = \text{sign } x$	162
12.4	$\tilde{H}^s$ -error for the one-dimensional problem with discontinuous right-hand side.	163
12.5	Adaptively refines meshes for the two-dimensional problem with constant right-hand side.	165
12.6	$\tilde{H}^s$ - and L^s -error for the two-dimensional problem with constant right-hand side.	166
12.7	Timings for assembly of the stiffness matrix, solution of linear system and computation of the error indicators for the two-dimensional problem.	167
12.8	Solutions u corresponding to the discontinuous right-hand side . .	167
12.9	Adaptively refined meshes for the discontinuous right-hand side. . .	168
12.10	$\tilde{H}^s$ -error for the two-dimensional problem with discontinuous right-hand side.	169
12.11	Solutions u on the L-shaped domain.	171
12.12	Adaptively refined meshes for the L-shaped domain.	171
12.13	$\tilde{H}^s$ -, L^2 -error and estimated error η for the L-shaped domain. . . .	172
12.14	Periodic solution $u = u_{2,8}^{2D}$	174
12.15	Timings in seconds for CG and MG depending on Δt	177
12.16	Number of iterations for CG and MG depending on Δt	178
12.17	Error in L^2 -norm and $\tilde{H}^s(\Omega)$ -norm in the Brusselator model. . . .	180
12.18	Localized spot solutions of the Brusselator system.	182

12.19 Stripe solutions of the Brusselator system.	183
---	-----

Part I

Fault-Prone Iterative Methods

Chapter 1

Multigrid Algorithm

Solvers for linear systems of equations can be classified into two categories. *Direct methods* aim to obtain the exact solution of the linear system (up to rounding error). This is typically achieved computing a factorization of the system matrix. Solution of the linear systems involving the factors is significantly easier, since their special structure can be exploited. The most well-known example of a direct method is Gauss elimination in which the system matrix is factored into upper and lower triangular factors. The drawback of direct methods is that even if the system matrix is sparse, i.e. if only $\mathcal{O}(n)$ of its entries are non-zeros, where n is the number of unknowns, the arising factors will be dense. Moreover, the computational complexity of direct methods is generally of order $\mathcal{O}(n^3)$. These drawbacks in memory usage and complexity motivated the development of *iterative methods*. These methods aim to obtain approximations that converge to the exact solution. The iterates are obtained only by matrix-vector products with the system matrix (and sometimes its transpose), which has complexity proportional to the number of non-zero entries of the matrix. In practice, obtaining only an approximate solution is not an issue. It suffices that the error in the solution of the system is dominated by other inherent errors, such as the one stemming from discretization of a continuum model. In the present chapter, we give an introduction to linear iterative methods, and to the

multigrid algorithm in particular.

1.1 Classical Iterative Solvers

Let $\mathbf{A}$ be a symmetric, positive definite sparse matrix arising from the discretization of an elliptic partial differential equation in d spatial dimensions, and $b \in \mathbb{R}^n$. The equation

$$\mathbf{A}x = b$$

can be rewritten as the fixed point problem

$$x = x + \mathbf{N}(b - \mathbf{A}x),$$

where $\mathbf{N}$ is a suitably chosen invertible matrix. If we require that the application of the so-called *preconditioner* $\mathbf{N}$ is inexpensive, evaluation of the right-hand side will not be costly because of the sparsity of $\mathbf{A}$. This leads to the following linear iterative method

$$x_{j+1} = x_j + \mathbf{N}(b - \mathbf{A}x_j).$$

If the method is to be used as a solver, then we need that $x_j \rightarrow x$ as $j \rightarrow \infty$. The error of the method after the $j + 1$ -th step can be written as

$$\begin{aligned} e_{j+1} &= x_{j+1} - x \\ &= x_j - x + \mathbf{N}(b - \mathbf{A}x_j) \\ &= (\mathbf{I} - \mathbf{N}\mathbf{A})e_j. \end{aligned}$$

We call $\mathbf{E} = \mathbf{I} - \mathbf{N}\mathbf{A}$ the *iteration matrix* or *error propagator*. The error converges to zero if and only if $\rho(\mathbf{E}) < 1$. The smaller the spectral radius of $\mathbf{E}$, the faster the convergence to the exact solution. The ideal preconditioner would be $\mathbf{A}^{-1}$, which would lead to convergence after just one iteration. Since finding $\mathbf{A}^{-1}$ is more

expensive than solving the original linear system $\mathbf{A}x = b$, $\mathbf{A}^{-1}$ is not a reasonable choice for the preconditioner. It demonstrates, however, that a good preconditioner should mimic the spectrum of $\mathbf{A}^{-1}$. On the other hand, taking $\mathbf{N} = \rho(\mathbf{A}) \mathbf{I}$ leads to a very inexpensive preconditioner, which, in turn, has often poor convergence properties since almost no information of the original problem is taken into account. The construction of efficient preconditioners that are both inexpensive and a good approximation to $\mathbf{A}^{-1}$ is an extensive field of very active research.

A classical iterative scheme is the (damped) Jacobi method, where $\mathbf{N}$ is proportional to the inverse of the diagonal part $\mathbf{D}$ of $\mathbf{A}$. The convergence rate is given by the following theorem:

Theorem 1.

Let $\mathbf{E} = \mathbf{I} - \theta \mathbf{D}^{-1} \mathbf{A}$ be the iteration matrix of the damped Jacobi method.

1. If $0 < \theta \mathbf{A} < 2\mathbf{D}$, then the damped Jacobi method converges.
2. If $0 < \lambda \mathbf{D} \leq \theta \mathbf{A} \leq \Lambda \mathbf{D}$, then the spectrum of $\mathbf{E}$ is contained in $[1 - \Lambda, 1 - \lambda]$ and the convergence rate is $\rho(\mathbf{E}) = \max\{\Lambda - 1, 1 - \lambda\}$.

Example 1.

Let $\mathbf{A}_{1D} \in \mathbb{R}^{n \times n}$ be the finite difference approximation to the 1D-Laplacian with Dirichlet boundary conditions:

$$\mathbf{A}_{1D} = \text{tridiag}_n(-1, 2, -1).$$

The eigenvalues of $\mathbf{A}_{1D}$ are

$$\lambda_k = 2 - 2 \cos \frac{k\pi}{n+1} \quad k = 1, \dots, n$$

and the corresponding eigenvectors are the discrete sine waves

$$\vec{v}_k = \left\{ \sin \frac{kj\pi}{n+1} \right\}_{j=1, \dots, n} \quad k = 1, \dots, n$$

Because $\mathbf{D} = 2\mathbf{I}$, we have $0 < \lambda\mathbf{D} \leq \mathbf{A}_{1D} \leq \Lambda\mathbf{D}$ with

$$\lambda = 1 - \cos\left(\frac{\pi}{n+1}\right),$$

$$\Lambda = 1 + \cos\left(\frac{\pi}{n+1}\right).$$

Hence

$$\rho(\mathbf{E}) = \cos\left(\frac{\pi}{n+1}\right).$$

For problem size n , we have

$$\rho(\mathbf{E}) = 1 - \frac{\pi^2}{2(n+1)^2} + \mathcal{O}((n+1)^{-4})$$

which makes this method quite slow. We observe though that the error components along the different eigenvectors $\vec{v}_k$ get damped by

$$\rho_k = 1 - \frac{\lambda_k}{2} = \cos\left(\frac{k\pi}{n+1}\right).$$

The components along $\vec{v}_k$, where k is close to $\frac{n+1}{2}$ decrease indeed much faster.

Example 1 shows that the Jacobi iteration reduces high-frequency error (with large λ_k) rapidly, but stalls on low frequency error. This suggests coupling the Jacobi iteration with a second type of iteration that specifically targets the low-frequency error.

1.2 Multigrid Algorithm

The fundamental observation that leads to the multigrid method is that low-frequency error varies slowly, whereas high frequency error is oscillatory. This means that low-frequency error can be transferred to a coarser discretization without losing too much information.

In what follows, let $\Omega \subset \mathbb{R}^d$ be a convex polyhedral domain, $\Gamma_D \subseteq \partial\Omega$ be non-empty and $V = \{v \in H^1(\Omega) : v = 0 \text{ on } \Gamma_D\}$. Let $\mathcal{T}_0$ be a partitioning of Ω into

the union of elements each of which is a $d + 1$ -simplex, such that the intersection of any pair of distinct elements is a single common vertex, edge, face or sub-simplex of both elements. For $\ell \in \mathbb{N}$, the partition $\mathcal{T}_\ell$ is obtained by uniformly subdividing each element in $\mathcal{T}_{\ell-1}$ into sub-simplices such that $\mathcal{T}_\ell$ satisfies the same conditions as $\mathcal{T}_0$ [73]. Let $\mathbb{P}_k$, $k \in \mathbb{N}$, consist of polynomials on $\mathbb{R}^d$ of total degree at most k . The subspace V_ℓ consists of continuous piecewise polynomials belonging to $\mathbb{P}_k$ on each element on the partition $\mathcal{T}_\ell$. In particular, this means that the spaces are nested: $V_0 \subset V_1 \subset \dots \subset V_L \subset V$.

Let $a : V \times V \rightarrow \mathbb{R}$ be a continuous, V -elliptic, bilinear form and $F : V \rightarrow \mathbb{R}$ be a continuous linear form and consider the problem of obtaining $u_L \in V_L$ such that

$$a(u_L, v) = F(v) \quad \forall v \in V_L. \quad (1.1)$$

Let $\phi_\ell^{(i)}$, $i = 1, \dots, n_\ell$, denote the usual nodal basis for V_ℓ , and ϕ_ℓ be the corresponding vector formed using the basis functions.

Problem eq. (1.1) is then equivalent to the matrix equation

$$\mathbf{A}_L x_L = F_L \quad (1.2)$$

where $\mathbf{A}_\ell = a(\phi_\ell, \phi_\ell)$ is the stiffness matrix, x_L is the coefficient vector of u_ℓ and $F_L = F(\phi_L)$ is the load vector. For future reference, we define the mass matrix by $\mathbf{M}_\ell = (\phi_\ell, \phi_\ell)$.

The multigrid method solves the linear system of equations

$$\mathbf{A}x = b$$

by introducing a hierarchy of coarser problems

$$\mathbf{A}_\ell x_\ell = b_\ell$$

of size n_ℓ , $\ell = 0, \dots, L$ with $\mathbf{A}_L = \mathbf{A}$. Information gets transferred between levels through *restriction* and *prolongation* operators

$$\mathbf{R}_\ell^{\ell+1} : \mathbb{R}^{n_{\ell+1}} \rightarrow \mathbb{R}^{n_\ell}, \quad \mathbf{P}_{\ell+1}^\ell : \mathbb{R}^{n_\ell} \rightarrow \mathbb{R}^{n_{\ell+1}}.$$

We will assume that $\mathbf{P}_{\ell+1}^\ell = (\mathbf{R}_\ell^{\ell+1})^T$ along with the *Galerkin relation*

$$\mathbf{A}_\ell = \mathbf{R}_\ell^{\ell+1} \mathbf{A}_{\ell+1} \mathbf{P}_{\ell+1}^\ell.$$

In the finite element context, the natural way to define these operators is to employ the nesting of the spaces $V_\ell \subset V_{\ell+1}$, which means that there exists a matrix $\mathbf{R}_{\ell+1}^\ell$ such that $\phi_\ell = \mathbf{R}_{\ell+1}^\ell \phi_{\ell+1}$. We will drop sub- and superscripts on restriction and prolongation operators in what follows.

Moreover, linear iterations called *smootheners* are defined on each level as

$$x_\ell \leftarrow x_\ell + \mathbf{N}_\ell (b_\ell - \mathbf{A}_\ell x_\ell),$$

where $\mathbf{N}_\ell$ is a suitable preconditioner, e.g. the damped Jacobi preconditioner $\mathbf{N}_\ell = \theta \mathbf{D}_\ell^{-1}$, with $\mathbf{D}_\ell$ the diagonal part of $\mathbf{A}_\ell$, as described above.

The multigrid method now consists in applying several steps of smoothening, i.e. application of an iterative methods that reduces high frequency error, followed by restriction of the *residual* $b_\ell - \mathbf{A}_\ell x_\ell$ to the next coarser grid. The coarser problem is solved by one or more recursive calls to the multigrid algorithm, or, in case of the coarsest mesh, by a direct solver. The multigrid method is given in Algorithm 1 and illustrated in Figure 1.1 and may be invoked to obtain an approximate solution of problem eq. (1.2) by a call to $MG_L(b, 0)$. The reason why the smoothening step appears twice is that it can be beneficial for the iteration to be symmetric, since the corresponding preconditioner will then be symmetric. This permits, for example, to use a multigrid iteration as a preconditioner in the Conjugate Gradient method. The number of recursive calls γ to the multigrid method on a coarse grid leads to different multigrid schemes. The most popular ones are called the *V-cycle* for which $\gamma = 1$, and the *W-cycle* with $\gamma = 2$.

Alternative choices for the smoothening iteration are possible. Popular options apart from the damped Jacobi iteration include Gauss-Seidel iterations, SOR smootheners, block-smootheners and polynomial smootheners.

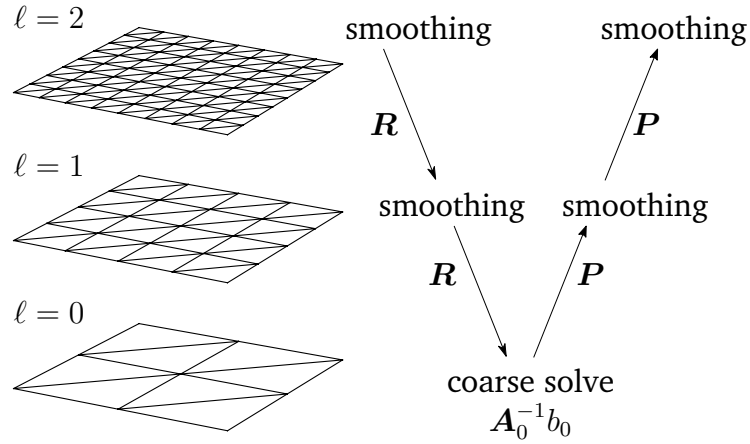


Figure 1.1: Multigrid hierarchy with three levels, and restrictions R and prolongations P .

Algorithm 1 Multigrid method MG_ℓ

Input: Right hand side b_ℓ ; Initial iterate x_ℓ

Output: $MG_\ell(b_\ell, x_\ell)$

- 1: **if** $\ell = 0$ **then return** $A_0^{-1}b_0$ // Exact solve on coarsest grid
 - 2: **else**
 - 3: **for** $i \leftarrow 1$ **to** ν_1 **do**
 - 4: $x_\ell \leftarrow S_\ell(b_\ell, x_\ell)$ // ν_1 pre-smoothing steps
 - 5: $d_{\ell-1} \leftarrow R(b_\ell - A_\ell x_\ell)$ // Restriction to coarser grid
 - 6: $e_{\ell-1} \leftarrow 0$
 - 7: **for** $j \leftarrow 1$ **to** γ **do**
 - 8: $e_{\ell-1} \leftarrow MG_{\ell-1}(d_{\ell-1}, e_{\ell-1})$ // γ coarse grid correction steps
 - 9: $x_\ell \leftarrow x_\ell + P e_{\ell-1}$ // Prolongation to finer grid
 - 10: **for** $i \leftarrow 1$ **to** ν_2 **do**
 - 11: $x_\ell \leftarrow S_\ell(b_\ell, x_\ell)$ // ν_2 post-smoothing steps
-

What makes the multigrid method particularly attractive is its computational complexity. On every levels, only sparse operations with complexity proportional to n_ℓ are used. Provided that the coarsening ratio from level to level is large enough, a single iteration of the multigrid method can easily be shown to have complexity $\mathcal{O}(n_L)$ or $\mathcal{O}(n_L \log n_L)$ [62]. As we will discuss in the next section, the rate of convergence of the multigrid method does not depend on the problem size. This means that an essentially constant amount of iterations needs to be taken to converge to a given error tolerance, making the multigrid method a solver with linear complexity.

1.3 Convergence Theory for Multigrid

The analysis of the convergence behavior of multigrid algorithms is a rich topic. Two theoretical frameworks are available. The first one relies on splitting the problem into high and low frequencies for a two grid method. The results are then extended to the multigrid case by a perturbation argument and recursive inequalities [62, 63]. The second framework is inspired by the analysis of domain decomposition methods and relies on a nested decomposition of the solution space [21, 101, 102]. In what follows, we adhere to the first framework as it will be seen to extend to the analysis of fault-prone methods.

Let the iteration matrix of pre- and post-smoother be given by E_ℓ^S . Then the iteration matrix of the multigrid method is given by

$$\begin{aligned} E_0 &= 0, \\ E_\ell &= (E_\ell^S)^{\nu_2} E_\ell^{CG} (E_\ell^S)^{\nu_1}, \end{aligned} \tag{1.3}$$

where the coarse grid correction is given by

$$E_\ell^{CG} = I - P (I - E_{\ell-1}^\gamma) A_{\ell-1}^{-1} R A_\ell.$$

In the case of a two grid method, the iteration matrix is given by

$$E_\ell^{TG} = (E_\ell^S)^{\nu_2} (I - P A_{\ell-1}^{-1} R A_\ell) (E_\ell^S)^{\nu_1}. \tag{1.4}$$

The standard assumptions for the convergence analysis of the multigrid method read as follows [20, 62, 63, 95]:

Assumption 1 (Smoothing property): There exists $\eta : \mathbb{N} \rightarrow \mathbb{R}_{\geq 0}$ satisfying $\lim_{\nu \rightarrow \infty} \eta(\nu) = 0$ and such that

$$\|\mathbf{A}_\ell (\mathbf{E}_\ell^S)^\nu\|_2 \leq \eta(\nu) \|\mathbf{A}_\ell\|_2, \quad \nu \geq 0, \ell = 1, \dots, L.$$

Assumption 2 (Approximation property): There exists a constant C_A such that

$$\|\mathbf{E}_\ell^{CG} \mathbf{A}_\ell^{-1}\|_2 \leq \frac{C_A}{\|\mathbf{A}_\ell\|_2}, \quad \ell = 1, \dots, L.$$

The smoothing and approximation property imply two grid convergence, with convergence rate given by

$$\begin{aligned} \rho(\mathbf{E}_\ell^{TG}) &\leq \left\| \mathbf{E}_\ell^{CG} (\mathbf{E}_\ell^S)^{\nu_1 + \nu_2} \right\|_2 \\ &\leq \left\| \mathbf{E}_\ell^{CG} \mathbf{A}_\ell^{-1} \right\|_2 \left\| \mathbf{A}_\ell (\mathbf{E}_\ell^S)^{\nu_1 + \nu_2} \right\|_2 \\ &\leq C_A \eta(\nu_1 + \nu_2). \end{aligned} \tag{1.5}$$

The right-hand side is less than one provided $\nu_1 + \nu_2$ is large enough.

Assumption 3: The smoother is non-expansive, i.e. $\rho(\mathbf{E}_\ell^S) = \|\mathbf{E}_\ell^S\|_A \leq 1$, and there exists $C_S > 0$ such that

$$\|(\mathbf{E}_\ell^S)^\nu\|_2 \leq C_S \quad \nu \geq 1, \ell = 1, \dots, L.$$

Assumption 3 permits to show that the two grid method is also a contraction with respect to the spectral norm $\|\bullet\|_2$:

$$\|\mathbf{E}_\ell^{TG}\|_2 \leq \|(\mathbf{E}_\ell^S)^{\nu_2}\|_2 \|\mathbf{E}_\ell^{CG} (\mathbf{E}_\ell^S)^{\nu_1}\|_2 \leq C_S C_A \eta(\nu_1).$$

While this result is weaker than eq. (1.5), it will be useful in the fault-prone case.

Assumption 4: There exist positive constants $\underline{C}_p$ and $\overline{C}_p$ such that

$$\underline{C}_p^{-1} \|x_\ell\|_2 \leq \|\mathbf{P}x_\ell\|_2 \leq \overline{C}_p \|x_\ell\|_2 \quad \forall x_\ell \in \mathbb{R}^{n_\ell}, \ell = 0, \dots, L-1.$$

Assumptions 3 and 4 allow to extend the convergence theory to the multilevel case with $\gamma \geq 2$. The most interesting case in practice is the W-cycle ($\gamma = 2$).

$$\begin{aligned}
\|\mathbf{E}_\ell\|_2 &= \|(\mathbf{E}_\ell^S)^{\nu_2} [\mathbf{I} - \mathbf{P}(\mathbf{I} - \mathbf{E}_{\ell-1}^\gamma) \mathbf{A}_{\ell-1}^{-1} \mathbf{R} \mathbf{A}_\ell] (\mathbf{E}_\ell^S)^{\nu_1}\|_2 \\
&\leq \|(\mathbf{E}_\ell^S)^{\nu_2} [\mathbf{I} - \mathbf{P} \mathbf{A}_{\ell-1}^{-1} \mathbf{R} \mathbf{A}_\ell] (\mathbf{E}_\ell^S)^{\nu_1}\|_2 \\
&\quad + \|(\mathbf{E}_\ell^S)^{\nu_2}\|_2 \|\mathbf{P}\|_2 \|\mathbf{E}_{\ell-1}\|_2^\gamma \|\mathbf{A}_{\ell-1}^{-1} \mathbf{R} \mathbf{A}_\ell (\mathbf{E}_\ell^S)^{\nu_1}\|_2 \\
&\leq \|\mathbf{E}_\ell^{TG}\|_2 + \underline{C}_p \overline{C}_p C_S \|\mathbf{E}_{\ell-1}\|_2^\gamma \|\mathbf{P} \mathbf{A}_{\ell-1}^{-1} \mathbf{R} \mathbf{A}_\ell (\mathbf{E}_\ell^S)^{\nu_1}\|_2.
\end{aligned}$$

Since

$$\begin{aligned}
\|\mathbf{P} \mathbf{A}_{\ell-1}^{-1} \mathbf{R} \mathbf{A}_\ell (\mathbf{E}_\ell^S)^{\nu_1}\|_2 &\leq \|(\mathbf{I} - \mathbf{P} \mathbf{A}_{\ell-1}^{-1} \mathbf{R} \mathbf{A}_\ell) (\mathbf{E}_\ell^S)^{\nu_1}\|_2 + \|(\mathbf{E}_\ell^S)^{\nu_1}\|_2 \\
&\leq C_S + C_A \eta(\nu_1),
\end{aligned}$$

we obtain the recursive inequality

$$\|\mathbf{E}_\ell\|_2 \leq \|\mathbf{E}_\ell^{TG}\|_2 + C_* \|\mathbf{E}_{\ell-1}\|_2^\gamma \quad (1.6)$$

with $C_* = C_S \underline{C}_p \overline{C}_p (C_S + C_A \eta(\nu_1))$. A classical result [63] concerning this inequality is

Lemma 2.

Suppose the elements of the sequence $\{\eta_\ell\}_{\ell \geq 1}$ satisfy $0 \leq \eta_1 \leq \xi$ and $0 \leq \eta_\ell \leq \xi + C_ \eta_{\ell-1}^\gamma$, $\ell \geq 2$. If $\gamma \geq 2$, $C_* \gamma > 1$ and $\xi \leq \frac{\gamma-1}{\gamma} \frac{1}{(C_* \gamma)^{\frac{1}{\gamma-1}}}$, then*

$$\eta_\ell \leq \begin{cases} \frac{\gamma}{\gamma-1} \xi, & \gamma \geq 2, \\ \frac{2\xi}{1+\sqrt{1-4C_*\xi}}, & \gamma = 2 \end{cases} < 1.$$

The result show that two grid convergence along with a sufficient number of smoothening steps imply that the multilevel scheme is convergent.

Chapter 2

Fault-Prone Iterative Methods

The taxonomy of faults and failures by Avižienis et al. [11] distinguishes between *hard faults* and *soft faults*. Soft faults correspond to corruption of instructions as well as the data produced and used within the code but which allow the execution of the code to proceed albeit in a corrupted state. An example of such a fault would be flipping of individual bits in a floating point number. Some bit flips in the exponent, sign or in the significant components of the mantissa may result in a relatively large corruption that becomes detectable. Hard faults might correspond to a message from one or more compute nodes either being corrupted beyond all recognition (possibly resulting in an exception being thrown), or lost altogether in the event of a compute node failing completely. Hard faults result in the interruption of the execution of the actual code unless remedial action is taken. As such, hard faults constitute errors that are readily *detectable*. Although the likelihood ε of a hard fault occurring can be assumed small, it is not negligible. The fact that an algorithm cannot continue after a hard fault without some kind of remedial action means that such faults cannot simply be ignored in the hope that the algorithm will recover naturally.

Most, if not all, existing algorithms were derived under the assumption that such faults cannot occur and accordingly, no indication is generally offered as to

what might constitute an appropriate course of action should a fault be detected in a particular algorithm. Nevertheless, in the event of a hard fault occurring, a decision has to be made as to what remedial action should be taken in order to resume the execution of the algorithm. The action that is chosen can have a dramatic effect on the performance and characteristics of the algorithm and as such represents a vital part of the algorithm. Ideally, the resulting algorithm should be subjected to the same kind of mathematical analysis that was applied to the original, deterministic, variant of the algorithm. Historical considerations mean that such analysis is available for very few methods.

At the present time, the basic approach to dealing with faults consists of attempting to restore the algorithm to the state that would have existed had no fault occurred. Two common ways in which this might be accomplished include: *Checkpoint-restart* [56, 64, 89] whereby the state of the system is stored (or checkpointed) at pre-determined intervals so that in the event of a fault occurring, the computation can be restarted from the stored state; or, *Process Replication* [48, 53, 64, 89], whereby each critical component of the overall process is replicated one or more times so that in the event of a component being subject to a fault, the true state can be restored by making reference to an unaffected replica. Hybrid variants of these strategies can also be contemplated. The drawbacks of such approaches are self-evident.

Different techniques have been previously employed in order to achieve fault resilience for iterative methods in general and multigrid in particular. Replication was used in [30, 38], and checkpoint-restart in [3, 4, 26]. Stoyanov and Webster [90] proposed a method based on selective reliability for fixed point methods, a concept that has been introduced in [46, 47, 65]. Finally, different authors [59, 69, 70] have proposed forward recovery methods to mitigate the effect of faults. The efficient detection of faults has been addressed in the context of various applications [14, 26, 81].

2.1 Modeling Faults

In order to take account of the effect of detectable faults on the stability and convergence of a numerical algorithm, we develop a simple probabilistic model for this behavior and describe how it is incorporated into the formulation of iterative algorithms. As a matter of fact, it will transpire that our model for detectable faults will also apply to the case of silent faults that are relatively small.

Our preferred approach to mitigating the effects of faults in multigrid and related iterative algorithms is rather simple: *the lost or corrupted result expected from a node is replaced by a zero*. That is to say, if a value $x \in \mathbb{R}$ is prone to possible hard faults, then we propose to continue the global computation using $\tilde{x} \in \mathbb{R}$ in place of x , where

$$\tilde{x} = \begin{cases} 0 & \text{if a fault is detected,} \\ x & \text{otherwise.} \end{cases}$$

This minimalist, *laissez-faire*, approach has obvious attractions in terms of restoring the system to a valid state without the need to (i) halt the execution (other than to perhaps migrate data to a substitute node in the event of a node failure); (ii) take or read any data checkpoints; (iii) recompute any lost quantities; or, (iv) compromise on resources through having to replicate processes. However, the efficacy of *laissez-faire* depends on the extent to which the convergence characteristics of the resulting algorithm mirror those of the fault-free deterministic variant.

2.2 Probabilistic Model of Faults

In order to model the effects of the *laissez-faire* fault mitigation approach on an algorithm, we introduce a Bernoulli random variable given by

$$\chi = \begin{cases} 0 & \text{with probability } q, \\ 1 & \text{with probability } 1 - q. \end{cases} \quad (2.1)$$

If a scalar variable x is subject to faults, then the effect of the fault and the laissez-faire strategy is modeled by replacing the value of x by the new value $\tilde{x} = \chi x$. Evidently $\tilde{x}$ is a random variable with mean and variance given by $\mathbb{E}[\tilde{x}] = (1 - q)x$ and $\mathbb{V}[\tilde{x}] = q(1 - q)x^2$. By the same token, the effect of a fault on a vector-valued variable $x \in \mathbb{R}^n$ is modeled in a similar fashion by defining $\tilde{x} = \mathcal{X}x$, where

$$\mathcal{X} = \text{diag}(\chi_1, \dots, \chi_n)$$

and χ_i are identically distributed Bernoulli random variables. The variables χ_i can be independent, thus modeling componentwise faults, or block-dependent, as would be the case for a node failure.

More formally, for given $\varepsilon > 0$, let $\mathbb{S}_\varepsilon$ denote a set consisting of random matrices satisfying the following conditions:

Assumption 5:

1. Each $\mathcal{X} \in \mathbb{S}_\varepsilon$ is a random diagonal matrix.
2. For every $\mathcal{X} \in \mathbb{S}_\varepsilon$, there holds $\mathbb{E}[\mathcal{X}] = e(\mathcal{X})\mathbf{I}$, where $e(\mathcal{X}) > 0$, and $|e(\mathcal{X}) - 1| \leq C\varepsilon$ for some fixed $C > 0$.
3. For every $\mathcal{X} \in \mathbb{S}_\varepsilon$ there holds $\|\mathbb{V}[\mathcal{X}]\|_2 = \max_{i,j} |\text{Cov}(\mathcal{X}_{ii}, \mathcal{X}_{jj})| \leq \varepsilon$.

The parameter ε is regarded as being small meaning that the random matrices $\mathcal{X}$ behave, with high probability, like an identity matrix. While the value of ε could vary from operation to operation, for example to take into account that denser matrix-vector products take more time and are therefore more likely to be hit by faults, we neglect this aspect in the analysis for the sake of clarity.

Important examples of random faults covered by Assumption 5 are

1. *Componentwise detectable faults*

A typical example of a detectable fault would be flipping of individual bits in a floating point number, resulting in a large enough upset to be detected,

or a pointer corruption that leads to an invalid memory address. Such cases can be treated using the laissez-faire strategy described above, and therefore modeled as componentwise faults with

$$\mathcal{X} = \text{diag}(\chi_1, \dots, \chi_n), \quad (2.2)$$

where χ_i are independent and identically ε -distributed Bernoulli random variables, i. e.

$$\chi_i = \begin{cases} 0 & \text{with probability } \varepsilon, \\ 1 & \text{with probability } 1 - \varepsilon. \end{cases}$$

2. Blockwise detectable faults

In the event of a node failure and application of the laissez-faire strategy, all the components of a vector that were residing on the node will be zeroed out. This can be modeled by a random matrix with independent diagonal *blocks*:

$$\mathcal{X} = \text{blockdiag}(\chi_1 \mathbf{I}_1, \dots, \chi_N \mathbf{I}_N), \quad (2.3)$$

where χ_i are independent and identically ε -distributed Bernoulli random variables and $\mathbf{I}_j$ are identity matrices.

3. Silent faults

Some bit flips in the exponent, sign or in the significant components of the mantissa may result in a relatively large error that becomes detectable. Such cases could be treated using the laissez-faire strategy described earlier. However, in other cases such faults may give a relatively small error and, as a result, be difficult (or impossible) to identify. Such silent faults can be modeled by a random matrix

$$\mathcal{X} = \mathbf{I} + \text{diag}(\eta_1 \chi_1, \dots, \eta_n \chi_n) \quad (2.4)$$

where η_i are independent and identically distributed random variables, and χ_i are independent and identically distributed Bernoulli random variables, such that $\mathcal{X} \in \mathbb{S}_\varepsilon$. In particular, this includes the cases of frequent faults with small impact and of rare but large upsets.

The matrices $\mathcal{X}$ will appear at various points in our model for the Fault-Prone Multigrid Algorithm. It is easy to see that if $\mathcal{X}_1$ and $\mathcal{X}_2$ are independent diagonal matrices, then $\mathcal{X}_1\mathcal{X}_2$ is again a random diagonal matrix.

Sometimes, a random matrix will appear in two (or more) different places in a single equation and it will be important to clearly distinguish between the cases where (i) each of the two matrices is a *different* realization of the random matrix, or (ii) the matrices both correspond to the *same* realization of the random matrix. We shall adopt the convention whereby should a symbol appear twice, then (ii) holds: i.e. the two occurrences represent the *same realization* of the random matrix and the matrices are therefore identical. However, if the square or higher power of a random matrix appears then (i) holds: i.e. each of the matrices in the product is a *different* realization of the same random matrix.

2.3 Scope of Faults Covered by the Analysis

Our analysis will cover faults that can be represented by a random matrix belonging to $\mathbb{S}_\varepsilon$. This means that the analysis will cover each of the cases (i) when the fault is detectable and mitigated using the laissez-faire strategy, and (ii) when the fault is silent but is relatively small in the sense that it may be modeled using (2.4). However, our analysis will not cover the important case involving faults which result in entries in the matrices or right hand sides in the problem being corrupted. Equally well, our analysis will not cover the case of bit flips that result in a large relative error but which nevertheless remain undetected. Finally, this work is restricted to the solve phase of multigrid; we assume that the setup phase is protected from faults.

2.4 A Model for Fault-Prone Multigrid

The multigrid algorithm comprises a number of steps each of which may be affected by faults were the algorithm to be implemented on a fault-prone machine. In the absence of faults, a single iteration of multigrid replaces the current iterate x_ℓ by $MG_\ell(b_\ell, x_\ell)$ defined in Algorithm 1. However, in the presence of faults, a single iteration of multigrid means that x_ℓ is replaced by $\mathcal{MG}_\ell(b_\ell, x_\ell)$ where $\mathcal{MG}_\ell$ differs from MG_ℓ in general due to corruption of intermediate steps in the computations arising from faults. The object of this section is to develop a model for faults and define the corresponding operators $\mathcal{MG}_\ell$, $\ell \in \mathbb{N}$. For simplicity, we assume that the coarse grid is of moderate size meaning that the exact solve $\mathbf{A}_0^{-1}$ on the coarsest grid is not prone to faults, i.e. $\mathcal{MG}_0(d_0, \cdot) = \mathbf{A}_0^{-1}d_0$. This is a reasonable assumption when the size of the coarse grid problem is sufficiently small that either faults are not an issue or, if they are, then replication can be used to achieve fault resilience. The extension of the analysis to the case of fault-prone coarse grid correction does not pose any fundamental difficulties but is not pursued in the present work.

2.5 Faults in the Smoother

A single application of the smoother S_ℓ to an iterate x_ℓ takes the form

$$S_\ell(b_\ell, x_\ell) = x_\ell + \mathbf{N}_\ell(b_\ell - \mathbf{A}_\ell x_\ell) \quad (2.5)$$

with each of the interior sub-steps in eq. (2.5) being susceptible to faults.

The innermost step is the computation of the residual $\rho_\ell = b_\ell - \mathbf{A}_\ell x_\ell$. The action of the matrix $\mathbf{A}_\ell$ is applied repeatedly throughout the solution phase. We shall assume that neither the entries or structure of matrix $\mathbf{A}_\ell$ nor the right hand side b_ℓ are subject to corruption, only computation involving them. This would be the case were the information needed to compute the action of $\mathbf{A}_\ell$ placed in non-volatile random access memory (NVRAM) with relatively modest overhead. The net effect

$$\begin{aligned}
\text{Global:} \quad & x_\ell \longrightarrow x_\ell + \mathcal{X}_\ell^{(S)} \mathbf{N}_\ell (b_\ell - \mathbf{A}_\ell x_\ell) \\
\text{Distributed:} \quad & \begin{array}{l}
\cancel{\text{node 1}} \quad [x_\ell]_1 \longrightarrow [x_\ell]_1 + \cancel{[\mathbf{N}_\ell]_1 [b_\ell - \mathbf{A}_\ell x_\ell]_1} = [x_\ell]_1 \\
\text{node 2} \quad [x_\ell]_2 \longrightarrow [x_\ell]_2 + [\mathbf{N}_\ell]_2 [b_\ell - \mathbf{A}_\ell x_\ell]_2 \\
\text{node 3} \quad [x_\ell]_3 \longrightarrow [x_\ell]_3 + [\mathbf{N}_\ell]_3 [b_\ell - \mathbf{A}_\ell x_\ell]_3
\end{array}
\end{aligned}$$

Figure 2.1: Schematic representation of a node failure during smoothening and the remedial action taken by the laissez-faire approach in the case of three compute nodes. $[\bullet]_i$ represents the part of the quantity $\bullet$ local to node i .

of faults in the computation of ρ_ℓ is that the preconditioner $\mathbf{N}_\ell$ does not act on the true residual but rather on a corrupted version modeled by $\mathcal{X}_1 \rho_\ell$ where $\mathcal{X}_1$ is a random diagonal matrix.

By the same token, the action of $\mathbf{N}_\ell$ is prone to faults, meaning that the true result $\mathbf{N}_\ell \mathcal{X}_1 \rho_\ell$ may be corrupted and is therefore modeled by $\mathcal{X}_2 \mathbf{N}_\ell \mathcal{X}_1 \rho_\ell$, where $\mathcal{X}_2$ is yet another random diagonal matrix. The matrix $\mathbf{N}_\ell$ corresponding to damped Jacobi is diagonal and hence $\mathcal{X}_2 \mathbf{N}_\ell \mathcal{X}_1 = \mathcal{X}^{(S)} \mathbf{N}_\ell$, where $\mathcal{X}^{(S)} = \mathcal{X}_1 \mathcal{X}_2$ is again a random diagonal matrix. Consequently, the combined effect of the two sources of error can be modeled by a single random diagonal matrix.

In summary, our model for the action of a smoothener prone to faults consists of replacing the true smoothener S used in the pre- and post-smoothening step in the multigrid algorithm by the *non-deterministic smoothener*

$$\mathcal{S}_\ell(b_\ell, x_\ell) = x_\ell + \mathcal{X}_\ell^{(S)} \mathbf{N}_\ell (b_\ell - \mathbf{A}_\ell x_\ell) \quad (2.6)$$

in which $\mathcal{X}_\ell^{(S)}$ is a random diagonal matrix which models the effect of the random faulty nature of the underlying hardware. The model eq. (2.6) tacitly assumes that the current iterate x_ℓ remains fault-free. We illustrate the remedial action to a node failure in Figure 2.1.

2.6 Faults in Restriction, Prolongation and Coarse Grid Correction

The restriction of the residual described by the step

$$d_{\ell-1} = \mathbf{R}(b_\ell - \mathbf{A}_\ell x_\ell) \quad (2.7)$$

is prone to faults. Firstly, as in the case of the smootheners, the true residual ρ_ℓ is prone to corruption and is modeled by $\mathcal{X}_\ell^{(\rho)} \rho_\ell$. The resulting residual is then operated on by the restriction $\mathbf{R}$, which is itself prone to faults modeled using a random diagonal matrix $\mathcal{X}_{\ell-1}^{(R)}$. We arrive at the following model for the effect of faults on eq. (2.7):

$$d_{\ell-1} = \mathcal{X}_{\ell-1}^{(R)} \mathbf{R} \mathcal{X}_\ell^{(\rho)} (b_\ell - \mathbf{A}_\ell x_\ell).$$

The coarse grid correction $e_{\ell-1}$ is obtained by performing γ iterations of $\mathcal{MG}_{\ell-1}$ with data $d_{\ell-1}$ and a zero initial iterate. The effect of faults when applying the prolongation to $e_{\ell-1}$ is also modeled by a random diagonal matrix $\mathcal{X}_\ell^{(P)}$ leading to the following model for the effect of faults on the coarse grid correction and prolongation steps:

$$x_\ell \leftarrow x_\ell + \mathcal{X}_\ell^{(P)} P e_{\ell-1}.$$

2.7 Model for Multigrid Algorithm in Presence of Faults

Replacing each of the steps in the fault-free Multigrid Algorithm 1 with their non-deterministic equivalent yields the following model for the Fault-Prone Multigrid Algorithm 2 and defines the associated fault-prone multilevel operators $\mathcal{MG}_\ell$, $\ell \in \mathbb{N}$.

The classical approach to the analysis of iterative solution methods for linear systems uses the notion of an iteration matrix as given in 1.3. For example, in the case

Algorithm 2 Model for Fault-Prone Multigrid Algorithm $\mathcal{MG}_\ell$ where $\mathcal{X}_\ell^{(\bullet)}$ are random diagonal matrices.

Input: Right hand side b_ℓ ; Initial iterate x_ℓ

Output: $\mathcal{MG}_\ell(b_\ell, x_\ell)$

```

1: if  $\ell = 0$  then return  $\mathbf{A}_0^{-1}b_0$            // Exact solve on coarsest grid
2: else
3:   for  $i \leftarrow 1$  to  $\nu_1$  do
4:      $x_\ell \leftarrow \mathcal{S}_\ell(b_\ell, x_\ell)$            //  $\nu_1$  pre-smoothing steps
5:      $d_{\ell-1} \leftarrow \mathcal{X}_{\ell-1}^{(R)} \mathbf{R} \mathcal{X}_\ell^{(\rho)} (b_\ell - \mathbf{A}_\ell x_\ell)$  // Restriction to coarser grid
6:      $e_{\ell-1} \leftarrow 0$ 
7:     for  $j \leftarrow 1$  to  $\gamma$  do
8:        $e_{\ell-1} \leftarrow \mathcal{MG}_{\ell-1}(d_{\ell-1}, e_{\ell-1})$  //  $\gamma$  coarse grid correction steps
9:        $x_\ell \leftarrow x_\ell + \mathcal{X}_\ell^{(P)} \mathbf{P} e_{\ell-1}$  // Prolongation to finer grid
10:    for  $i \leftarrow 1$  to  $\nu_2$  do
11:       $x_\ell \leftarrow \mathcal{S}_\ell(b_\ell, x_\ell)$            //  $\nu_2$  post-smoothing steps

```

of the fault-free smoothing step $\mathbf{E}_\ell^S = \mathbf{I} - \mathbf{N}_\ell \mathbf{A}_\ell$ is the iteration matrix, with preconditioner $\mathbf{N}_\ell$. The corresponding result for the fault-prone smoother eq. (2.6) is given by

$$x - \mathcal{S}_\ell(\mathbf{A}_\ell x, y) = \mathcal{E}_\ell^S(x - y), \quad \forall x, y \in \mathbb{R}^{n_\ell} \quad (2.8)$$

where the iteration matrix $\mathcal{E}_\ell^S = \mathbf{I} - \mathcal{X}_\ell^{(S)} \mathbf{N}_\ell \mathbf{A}_\ell$ is now *random*. Similarly, we define the iteration matrix for the Fault-Prone Multigrid Algorithm 2 by the equation

$$x - \mathcal{MG}_\ell(\mathbf{A}_\ell x, y) = \mathcal{E}_\ell(x - y), \quad \forall x, y \in \mathbb{R}^{n_\ell}. \quad (2.9)$$

In particular, $x - \mathcal{MG}_0(\mathbf{A}_0 x, \bullet) = 0$ and hence $\mathcal{E}_0 = 0$, i.e. the zero matrix.

The ν_1 pre-smoothing steps in Algorithm 2 can be expressed in the form

$$x^{(0)} = x_\ell; \quad x^{(i)} = \mathcal{S}_\ell(b_\ell, x^{(i-1)}), \quad i = 1, \dots, \nu_1.$$

Using eq. (2.8) gives

$$x - x_\ell^{(i)} = \mathcal{E}_\ell^S(x - x_\ell^{(i-1)}), \quad i = 1, \dots, \nu_1$$

and hence

$$x - x_\ell^{(\nu_1)} = (\mathcal{E}_\ell^S)^{\nu_1} (x - x_\ell). \quad (2.10)$$

Using this notation means that the vector $d_{\ell-1}$ appearing in Algorithm 2 is given by

$$d_{\ell-1} = \mathcal{X}_{\ell-1}^{(R)} \mathbf{R} \mathcal{X}_{\ell}^{(\rho)} (b_{\ell} - \mathbf{A}_{\ell} x_{\ell}^{(\nu_1)}) = \mathcal{X}_{\ell-1}^{(R)} \mathbf{R} \mathcal{X}_{\ell}^{(\rho)} \mathbf{A}_{\ell} (x - x_{\ell}^{(\nu_1)}).$$

where $b_{\ell} = \mathbf{A}_{\ell} x$.

The coarse grid correction steps of Algorithm 2 can be written in the form

$$e_{\ell-1}^{(0)} = 0; \quad e_{\ell-1}^{(j)} = \mathcal{MG}_{\ell-1}(d_{\ell-1}, e_{\ell-1}^{(j-1)}), \quad j = 1, \dots, \gamma.$$

Let $z = \mathbf{A}_{\ell-1}^{-1} d_{\ell-1}$ then

$$z - e_{\ell-1}^{(j)} = z - \mathcal{MG}_{\ell-1}(d_{\ell-1}, e_{\ell-1}^{(j-1)}) = \mathcal{E}_{\ell-1}(z - e_{\ell-1}^{(j-1)}), \quad j = 1, \dots, \gamma$$

thanks to eq. (2.9). Iterating this result and recalling that $e_{\ell-1}^{(0)} = 0$, gives $z - e_{\ell-1}^{(\gamma)} = \mathcal{E}_{\ell-1}^{\gamma} z$, or, equally well,

$$e_{\ell-1}^{(\gamma)} = (\mathbf{I} - \mathcal{E}_{\ell-1}^{\gamma}) \mathbf{A}_{\ell-1}^{-1} d_{\ell-1}.$$

Using these notations, the prolongation to the finer grid step in Algorithm 2 takes the form

$$x_{\ell}^{(P)} = x_{\ell}^{(\nu_1)} + \mathcal{X}_{\ell}^{(P)} \mathbf{P} e_{\ell-1}^{(\gamma)}$$

and hence, by collecting results, we deduce that

$$\begin{aligned} x - x_{\ell}^{(P)} &= (x - x_{\ell}^{(\nu_1)}) - \mathcal{X}_{\ell}^{(P)} \mathbf{P} e_{\ell-1}^{(\gamma)} \\ &= \left[\mathbf{I} - \mathcal{X}_{\ell}^{(P)} \mathbf{P} (\mathbf{I} - \mathcal{E}_{\ell-1}^{\gamma}) \mathbf{A}_{\ell-1}^{-1} \mathcal{X}_{\ell-1}^{(R)} \mathbf{R} \mathcal{X}_{\ell}^{(\rho)} \mathbf{A}_{\ell} \right] (x - x_{\ell}^{(\nu_1)}). \end{aligned}$$

Arguments identical to those leading to eq. (2.10) give the following result

$$x - \mathcal{MG}_{\ell}(\mathbf{A}_{\ell} x, x_{\ell}) = (\mathcal{E}_{\ell}^S)^{\nu_2} (x - x_{\ell}^{(P)}).$$

which, in view of identity eq. (2.9) and eq. (2.10), yields the following recursive formula for the iteration matrix of the Fault-Prone Multigrid Algorithm 2:

$$\mathcal{E}_{\ell} = \left(\mathcal{E}_{\ell}^{S, \text{post}} \right)^{\nu_2} \left[\mathbf{I} - \mathcal{X}_{\ell}^{(P)} \mathbf{P} (\mathbf{I} - \mathcal{E}_{\ell-1}^{\gamma}) \mathbf{A}_{\ell-1}^{-1} \mathcal{X}_{\ell-1}^{(R)} \mathbf{R} \mathcal{X}_{\ell}^{(\rho)} \mathbf{A}_{\ell} \right] \left(\mathcal{E}_{\ell}^{S, \text{pre}} \right)^{\nu_1}, \quad (2.11)$$

Chapter 3

Lyapunov Exponents and Their Computation

In the previous chapter, we determined the random iteration matrices corresponding to the Fault-Prone Two Grid and Multigrid Methods. To analyze their respective convergence, we introduce the theory of products of random matrices that concerns itself with the study of the asymptotic behavior of product sequences of random matrices. Most importantly, we will state a “law of large numbers” for random matrices, and discuss a multitude of approaches to the computation and estimation of the asymptotic convergence rate.

3.1 Lyapunov Exponents

Let $\mathcal{A}$ be a random matrix of size $n \times n$ with distribution μ . We will assume that μ is a discrete probability distribution and picks values in a finite set $\mathcal{A} \subset GL_n$.

We define the *Lyapunov exponent* $\gamma(\mathcal{A})$ as

$$\gamma(\mathcal{A}) := \lim_{N \rightarrow \infty} \frac{1}{N} \mathbb{E} [\log \|\mathcal{A}^N\|] .$$

(Remember that according to our notation $\mathcal{A}^N = \prod_{j=1}^N \mathbf{A}_j$, where $\mathbf{A}_j$ are all independent and of the same distribution as $\mathcal{A}$.) This quantity does not depend on the choice of the norm $\|\cdot\|$ as all norms are equivalent in finite dimension. The definition is motivated by the equality

$$\log \rho(\mathcal{A}) = \lim_{N \rightarrow \infty} \frac{1}{N} \log \|\mathcal{A}^N\|$$

that holds for non-random matrices $\mathbf{A}$.

Furstenberg and Kesten[55] showed that provided that $\mathbb{E} [\log (\mathcal{A})^+] < \infty$, $\gamma(\mathcal{A})$ exists and

$$\gamma(\mathcal{A}) = \lim_{N \rightarrow \infty} \frac{1}{N} \log \|\mathcal{A}^N x_0\| \quad \text{a.s.}$$

The *Lyapunov spectral radius*

$$\varrho(\mathcal{A}) := e^{\gamma(\mathcal{A})} = \lim_{N \rightarrow \infty} \|\mathcal{A}^N x_0\|^{\frac{1}{N}} \quad \text{a.s.}$$

is the quantity corresponding to the spectral radius of a non-random matrix. $\varrho(\mathcal{A}) < 1$ means that the matrix product is convergent, whereas $\varrho(\mathcal{A}) > 1$ means that it is divergent.

We also define the *generalized Lyapunov exponents* $L(\mathcal{A}, \alpha)$ for $\alpha \neq 0$ as

$$L(\mathcal{A}, \alpha) := \lim_{N \rightarrow \infty} \frac{1}{N} \log \mathbb{E} [\|\mathcal{A}^N\|^\alpha]$$

and the *generalized Lyapunov spectral radii* $\varrho_\alpha(\mathcal{A})$ as

$$\varrho_\alpha(\mathcal{A}) := e^{\frac{1}{\alpha} L(\mathcal{A}, \alpha)}.$$

It can be shown that $\frac{1}{N} \log \|\mathcal{A}^N x_0\|$ satisfies a large deviation principle with a rate function I which means that

$$P\left(\frac{1}{N} \log \|\mathcal{A}^N x_0\| = \lambda\right) \approx e^{-NI(\lambda)}.$$

From the above, we find that I satisfies

$$I(\gamma(\mathcal{A})) = 0,$$

$$I'(\gamma(\mathcal{A})) = 0,$$

$$I''(\gamma(\mathcal{A})) > 0.$$

We can approximately calculate

$$\begin{aligned} \mathbb{E} [\|\mathcal{A}^N x_0\|^\alpha] &= \mathbb{E} \left[e^{N\alpha \frac{1}{N} \log \|\mathcal{A}^N x_0\|} \right] \\ &\approx \int e^{N\alpha\lambda} e^{-NI(\lambda)} d\lambda \\ &\approx e^{N \sup_\lambda \{\alpha\lambda - I(\lambda)\}} \end{aligned}$$

and hence

$$L(\mathcal{A}, \alpha) = \sup_\lambda \alpha\lambda - I(\lambda).$$

This means that L is the Fenchel-Legendre transform of the rate function I . For each α^* , there is a characteristic growth rate λ^* that corresponds to it, given by

$$\frac{\partial L}{\partial \alpha}(\alpha^*) = \lambda^*.$$

Therefore, the values of $L(\mathcal{A}, \alpha^*)$ depend on the unlikely sequences $\{\mathcal{A}^N x_0\}$ with growth rate λ^* different from $\gamma(\mathcal{A})$.

Finally, the Lyapunov spectral radius can also be estimated by neglecting the interactions between different values taken by the random matrix $\mathcal{A}$. We define

$$\bar{\varrho}(\mathcal{A}) := \exp \sup_{\|x\|=1} \mathbb{E} [\log \|\mathcal{A}x\|],$$

$$\tilde{\varrho}(\mathcal{A}) := \exp \mathbb{E} [\log \|\mathcal{A}\|]$$

and

$$\tilde{\varrho}_\alpha(\mathcal{A}) := \exp \frac{1}{\alpha} \log \mathbb{E} [\|\mathcal{A}\|^\alpha]$$

for induced sub-multiplicative norms. These quantities obviously depend on the choice of the norms.

Hence, we have for $\alpha > 0$

$$\begin{aligned}
\varrho_{-\alpha}(\mathcal{A}) &\leq \varrho(\mathcal{A}) \leq \varrho_{\alpha}(\mathcal{A}) \\
&\leq \bar{\varrho}(\mathcal{A}) \leq \tilde{\varrho}(\mathcal{A}) \\
&\leq \tilde{\varrho}_{-\alpha}(\mathcal{A}) \leq \tilde{\varrho}_{\alpha}(\mathcal{A})
\end{aligned}$$

The horizontal inequalities follow from Jensen's inequality.

Moreover, if μ picks only one matrix $\mathbf{A}$, so that we are in fact in the non-random case, we find

$$\varrho(\mathbf{A}) = \varrho_{\alpha}(\mathbf{A}) = \rho(\mathbf{A}).$$

The quantities $\bar{\varrho}(\mathbf{A})$, $\tilde{\varrho}(\mathbf{A})$ and $\tilde{\varrho}_{\alpha}(\mathbf{A})$ can then be made arbitrarily close to $\rho(\mathbf{A})$ by choosing suitable vector and matrix norms.

The theory of Lyapunov exponents is thoroughly discussed in the book [19] by Bougerol and Lacroix.

The following example adapted from [37] shows that the interaction between the values taken by $\mathcal{A}$ can be important:

Example 2.

Let $a \in \mathbb{R}$ and let the random matrix $\mathcal{A}$ take values

$$\begin{pmatrix} a & 0 \\ 0 & \frac{1}{2a} \end{pmatrix} \quad \text{with probability } (1-p) \quad \text{and} \quad \mathbf{S} := \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \quad \text{with probability } p.$$

We are interested in the growth and decay of

$$X_N := \mathcal{A}^N x_0 = \prod_{j=1}^N \mathbf{A}_j x_0,$$

where

$$x_0 = \alpha \vec{e}_1 + \beta \vec{e}_2$$

and $\mathbf{A}_j$ are independent random matrices that are distributed like $\mathcal{A}$. Obviously, we have $\varrho(\mathcal{A}) = \max\{a, \frac{1}{2a}\}$ for $p = 0$ and $\varrho(\mathcal{A}) = 1$ for $p = 1$.

We define the stopping times

$$\begin{aligned}\tau_0 &= 0, \\ \delta_k &= \inf\{i \geq 1 : \mathbf{A}_{\tau_{k-1}+i} = \mathbf{S}\}, & k \geq 1, \\ \tau_k &= \inf\{i \geq 1 : \mathbf{A}_{\delta_k+i} = \mathbf{S}\}, & k \geq 1.\end{aligned}$$

And set

$$\sigma_k = \sum_{n=1}^k \delta_n + \tau_n, \quad k \geq 1.$$

Between σ_k and $\sigma_k + \delta_{k+1}$ the first component is expanding by a and the second one is damped by $\frac{1}{2a}$. Between $\sigma_k + \delta_{k+1}$ and σ_{k+1} , the roles are exchanged. Moreover, δ_k, τ_k are all independent and identically geometrically distributed with parameter p . We have

$$\log |X_{j,1}| = \log |\alpha| + \log a \left[\sum_{n=1}^k (\delta_n - \tau_n) + N - \sigma_k \right] - \log 2 \sum_{n=1}^k (\tau_n - 1)$$

if $\sigma_k \leq N < \sigma_k + \delta_{k+1}$ or

$$\begin{aligned}\log |X_{j,1}| &= \log |\beta| + \log a \left[\sum_{n=1}^k (\tau_n - \delta_n) + 2\delta_{k+1} - N + \sigma_k - 1 \right] \\ &\quad - \log 2 \left[\sum_{n=1}^{k+1} (\tau_n - 1) + N - \sigma_k - \delta_{k+1} \right]\end{aligned}$$

if $\sigma_k + \delta_{k+1} \leq N < \sigma_{k+1}$. Similar expressions hold for $\log |X_{N,2}|$. Now,

$$\frac{\sigma_k}{k} \leq \frac{N}{k} \leq \frac{\sigma_{k+1}}{k}$$

and both upper and lower bound converges almost surely to $\mathbb{E}[\delta_1 + \tau_1] = \frac{2}{p}$.

$$\frac{1}{k} \sum_{n=1}^k \tau_n, \quad \frac{1}{k} \sum_{n=1}^k \delta_n \xrightarrow{a.s.} \mathbb{E}[\tau_1] = \frac{1}{p},$$

$$\frac{\delta_{k+1}}{N} \xrightarrow{a.s.} 0,$$

so we find

$$\frac{1}{N} \log |X_{j,1}| \xrightarrow{a.s.} -\log(2) \frac{p}{2} \left(\frac{1}{p} - 1 \right)$$

and

$$|X_{N,1}|^{\frac{1}{N}} \xrightarrow{a.s.} 2^{\frac{p-1}{2}}.$$

The same is true for the second component and hence

$$\varrho(\mathcal{A}) = \lim_{N \rightarrow \infty} \|X_N\|^{\frac{1}{N}} = \begin{cases} \max \left\{ a, \frac{1}{2a} \right\} & \text{for } p = 0, \\ 2^{\frac{p-1}{2}} & \text{for } 0 < p < 1, \\ 1 & \text{for } p = 1. \end{cases}$$

On the other hand,

$$\log \tilde{\varrho}(\mathcal{A}) = \mathbb{E} [\log \|\mathcal{A}\|_2] = (1-p) \log \max \left\{ a, \frac{1}{2a} \right\}$$

and

$$\log \bar{\varrho}(\mathcal{A}) = \sup_{\|x\|_2=1} \mathbb{E} [\log \|\mathcal{A}x\|_2] = (1-p) \log \max \left\{ a, \frac{1}{2a} \right\}.$$

So for $0 < p < 1$ and $a \notin [\frac{1}{2}, 1]$,

$$\varrho(\mathcal{A}) < 1 < \tilde{\varrho}(\mathcal{A}) = \bar{\varrho}(\mathcal{A}).$$

Moreover, $\tilde{\varrho}(\mathcal{A}), \bar{\varrho}(\mathcal{A}) \rightarrow \infty$ as $a \rightarrow \infty$ or $a \rightarrow 0$, whereas $\varrho(\mathcal{A})$ is independent of a .

Remark: The previous example shows that the Lyapunov spectral radius can be discontinuous with respect to the weights and matrices in the support of $\mathcal{A}$. In particular, this means that we cannot conclude from $\rho(\mathbf{E}_l) < 1$ that $\varrho(\mathcal{E}_l) < 1$ for small perturbations.

The next example shows that measuring unlikely sequences can lead to serious over-estimation of the convergence rate:

Example 3.

Let $\mathcal{A}$ be a random 1×1 matrix taking values 2^a with probability $\frac{1}{a}$ and 1 with probability $\frac{a-1}{a}$. Hence the growth rates of all possible sequences are between 1 and 2^a . Then

$$\varrho(\mathcal{A}) = \bar{\varrho}(\mathcal{A}) = \tilde{\varrho}(\mathcal{A}) = 2,$$

$$\varrho_\alpha(\mathcal{A}) = \left(\frac{1}{a} 2^{a\alpha} + \frac{a-1}{a} \right)^{\frac{1}{\alpha}},$$

and for large a we find $\varrho(\mathcal{A}) \ll \varrho_\alpha(\mathcal{A})$, $\alpha > 0$.

The following theorem is given as an exercise in the book [19].

Theorem 3.

Suppose that a random matrix A has block triangular form

$$\mathcal{A} = \begin{pmatrix} \mathcal{B} & \mathcal{C} \\ 0 & \mathcal{D} \end{pmatrix}$$

with square blocks $\mathcal{B}$ and $\mathcal{D}$. Then

$$\gamma(\mathcal{A}) = \max\{\gamma(\mathcal{B}), \gamma(\mathcal{D})\}.$$

An immediate consequence is

Lemma 4.

Suppose that the support $\mathcal{A}$ of a random matrix $\mathcal{A}$ can be simultaneously triangularized, i.e. there exists an invertible matrix P such that $P^{-1}BP$ is upper triangular for every $B \in \mathcal{A}$. Then

$$\gamma(\mathcal{A}) = \max_k \mathbb{E} \left[\log (P^{-1} \mathcal{A} P)_{kk} \right].$$

In general, unless the matrices in the support of $\mathcal{A}$ can be simultaneously triangularized, it is impossible to calculate $\varrho(\mathcal{A})$ analytically. Several different methods to do so numerically have been proposed:

- Monte-Carlo simulation of trajectories of $\mathcal{A}^N x_0$,
- Approximation of the invariant distribution of the Markov chain $\frac{\mathcal{A}^N x_0}{\|\mathcal{A}^N x_0\|}$,
- cycle expansions,
- weak disorder expansions.

A good review of these methods is given in [37]. We discuss and use the Monte Carlo approach as well as the calculation of the invariant distribution below. The cycle expansion can only be applied successfully if the cardinality of $\mathcal{A}$ is small, whereas the weak disorder expansion works only for small deviations around a deterministic matrix. As the invariant distribution method is limited to low dimensional matrices, the Monte-Carlo approach is the only generally available option to obtain the asymptotic growth rate $\varrho(\mathcal{A})$.

The situation is slightly better for the generalized Lyapunov spectral radius $\varrho_\alpha(\mathcal{A})$ which can be calculated using

1. Vanneste's [97] resampled Monte-Carlo algorithm,
2. the Replica Trick.

We shortly describe both methods below. The resampled Monte-Carlo algorithm can be seen as the equivalent of the Monte-Carlo approach for the calculation of $\varrho(\mathcal{A})$ in that it is generally available. The more specialized Replica Trick can only be used to calculate generalized Lyapunov exponents of positive even order.

3.2 Monte-Carlo Approximation

We can approximate the Lyapunov exponent $\gamma(\mathcal{A})$ by Monte-Carlo simulations of a trajectory

$$X_N(\omega) = \prod_{j=1}^N \mathbf{A}_j(\omega) x_0$$

for normalized random initial vector x_0 and calculating

$$\gamma(\mathcal{A}) \approx \frac{1}{N} \log \|X_N(\omega)\|.$$

This makes sense because of the almost sure convergence of X_N to $\gamma(\mathcal{A})$.

In order to avoid over- and underflow of the components of the vector $X_N(\omega)$, we renormalize after every step, i.e. we set

$$\begin{aligned} \tilde{X}_j &:= \frac{\mathbf{A}_j(\omega) \tilde{X}_{j-1}}{\|\mathbf{A}_j(\omega) \tilde{X}_{j-1}\|} \\ \alpha_j &:= \|\mathbf{A}_j(\omega) \tilde{X}_{j-1}\| \end{aligned}$$

and approximate

$$\gamma(\mathcal{A}) \approx \frac{1}{N} \sum_{j=1}^N \log \alpha_j.$$

3.3 Approximation via Invariant Distribution

Another characterization of γ can be given in terms of the invariant probability distribution ν of the Markov chain $\tilde{X}_N = \frac{X_N}{\|X_N\|}$ on the projective space $P\mathbb{R}^{d-1}$:

$$\gamma = \int_{P\mathbb{R}^{d-1}} d\nu(x) \int_{\mathbb{R}^{d \times d}} d\mu(A) \log \|Ax\|$$

Here, ν satisfies the invariance equation

$$\int_{P\mathbb{R}^{d-1}} d\nu(x) f(x) = \int_{P\mathbb{R}^{d-1}} d\nu(x) \int_{\mathbb{R}^{d \times d}} d\mu(A) f\left(\frac{Ax}{\|Ax\|}\right)$$

for any Borel measurable bounded function f on $P\mathbb{R}^{d-1}$.

If d is small and only finitely many matrices are picked by μ , it is possible to approximate ν and then calculate γ . The following method has been proposed by Embree and Trefethen in [49] for the calculation of the Lyapunov spectral radius of the random Fibonacci sequence

$$x_{n+1} = x_n \pm x_{n-1},$$

where the signs $=$ or $-$ are chosen with equal probability. Rewritten as a one-step sequence, this corresponds to picking the matrices

$$\mathcal{A} = \begin{pmatrix} 0 & 1 \\ \pm 1 & 1 \end{pmatrix}$$

with equal probability. This has become a test problem for different algorithms to determine the Lyapunov spectral radius, because Viswanath [98] was able to determine with high precision that $\varrho(\mathcal{A}) = 1.1319882\dots$

We restrict ourselves to 2×2 matrices. Vectors $\vec{x}$ of $P\mathbb{R}$ can be parametrized by their slope $m = \frac{x_2}{x_1} \in (-\infty, \infty) \cup \{\infty\}$ or by their angle $\theta = \arctan\left(\frac{x_2}{x_1}\right) \in (-\frac{\pi}{2}, \frac{\pi}{2}]$. The action of a matrix B on an angle θ or a slope m can be written as

$$\begin{aligned} B \cdot \theta &= \arctan \frac{y_2}{y_1}, & \vec{y} &= B \begin{pmatrix} \cos \theta \\ \sin \theta \end{pmatrix} \\ B \cdot m &= \frac{y_2}{y_1}, & \vec{y} &= B \begin{pmatrix} 1 \\ m \end{pmatrix} \end{aligned}$$

We call θ or m a point of discontinuity of B , if $y_1 = 0$.

For numerical calculations, it is easier to work with a bounded interval, so we use the second option and partition the angular domain into K intervals:

$$(-\frac{\pi}{2}, \frac{\pi}{2}] = \bigcup_{k=1}^K I_k.$$

We then make the following ansatz:

$$\begin{aligned}\nu(J) &= \sum_{k=1}^K \nu(I_k \cap J) \\ &\approx \sum_{k=1}^K \nu_k \frac{l(I_k \cap J)}{l(I_k)}\end{aligned}$$

where $l(J)$ is the length of the interval J . We plug the ansatz into the invariance equation to obtain the set of equations

$$\begin{aligned}\nu_m = \nu(I_m) &= \sum_{k=1}^K \nu_k \int d\mu(\mathbf{A}) \frac{l(I_k \cap \mathbf{A}I_m)}{l(I_k)} \\ &=: \sum_{k=1}^K \nu_k B_{m,k}.\end{aligned}$$

Here, we write $\mathbf{A}I_k$ to designate the image of the angular interval I_k under the action of $\mathbf{A}$.

Supplemented with the condition

$$1 = \sum_{k=1}^K \nu_k$$

there exists a unique solution vector ν_k of the resulting system

$$\begin{pmatrix} \mathbf{B} - \mathbf{I} \\ \mathbf{1} \end{pmatrix} \nu = \begin{pmatrix} \vec{0} \\ 1 \end{pmatrix}.$$

It can be solved by either replacing one of the first K equations by the last one and then solving a square system or as a least squares problem.

Now, γ can be approximated as

$$\gamma \approx \sum_{j=1}^K \nu_j \int d\mu(\mathbf{A}) \log \|\mathbf{A}x_j\|.$$

We chose the evaluation point $x_k \in I_k$ to be the midpoint.

It remains to pick a suitable partition $\mathcal{I} = \{I_k\}$ of the angular domain. It is convenient to choose the points of discontinuity for all matrices in $\mathcal{A}$ as boundaries for

the intervals in order to simplify the calculation of $l(I_k \cap AI_m)$. Embree and Trefethen [49] used uniformly spaced nodes for the Fibonacci problem. (This works, because the discontinuities are $\pm \frac{\pi}{4}$.) In order to get precision of 6 digits, they had to pick 2^{20} nodes, resulting in a matrix with 10^7 non-zero entries. Instead, we propose the following strategy of choosing a partition of $(-\frac{\pi}{2}, \frac{\pi}{2}]$:

Algorithm 4 Invariant distribution algorithm with refinement

Input: Number of refinement steps N , refinement threshold θ , number of subdivisions k

- 1: **for** $i \leftarrow 1$ **to** N **do**
 - 2: $\mathcal{I} \leftarrow \text{refine}(\mathcal{I}, \nu; \theta, k)$
 - 3: $\nu \leftarrow \text{FindInvariantDistribution}(\mathcal{I})$
 - 4: $\gamma \leftarrow \text{CalculateLyapunovExponent}(\mathcal{I}, \nu)$
-

This strategy does worst, when the invariant distribution puts similar weight on the whole angular domain, and best if it has localized support. The method can obviously only be applied to low dimensional problems, as it requires the discretization of the hyper-cube $(-\frac{\pi}{2}, \frac{\pi}{2})^{d-1}$.

3.4 Resampled Monte-Carlo Method

The generalized Lyapunov exponents $L(\alpha)$ could in principle be calculated just as the Lyapunov exponent γ using brute-force Monte Carlo simulations:

$$\begin{aligned} L(\alpha) &= \lim_{N \rightarrow \infty} \frac{1}{N} \log \mathbb{E} [\|\mathbf{X}_N\|^\alpha] \\ &\approx \frac{1}{N} \log \left[\frac{1}{K} \sum_{k=1}^K \left\| X_N^{(k)} \right\|^\alpha \right], \end{aligned}$$

where $X_N^{(k)}$ are independent realizations of $\mathbf{X}_N = \prod_{j=1}^N \mathcal{A}_j x_0$. Vanneste[97] remarked that because the generalized Lyapunov exponents are associated with a class of sequences that display a growth behavior different from γ , it would be advantageous to give those rare sequences a higher weight. He proposed a resampling

strategy that selects exactly those realizations. At every step of the simulation, we calculate as before

$$\begin{aligned}\bar{X}_j^{(k)} &:= \frac{\mathbf{A}_j^{(k)} \tilde{X}_{j-1}^{(k)}}{\left\| \mathbf{A}_j^{(k)} \tilde{X}_{j-1}^{(k)} \right\|}, \\ \alpha_j^{(k)} &:= \left\| \mathbf{A}_j^{(k)} \tilde{X}_{j-1}^{(k)} \right\|^\alpha.\end{aligned}$$

But then we resample, i.e. we set

$$\beta_j = \sum_{k=1}^K \alpha_j^{(k)}$$

and pick J_k for $k = 1, \dots, K$ in $\{1, \dots, K\}$ with probabilities $\frac{\alpha_j^{(k)}}{\beta_j}$. Then we assign

$$\tilde{X}_j^{(k)} = \bar{X}_{J_k}^{(k)}.$$

Finally, we approximate

$$L(\alpha) \approx \frac{1}{N} \sum_{j=1}^N \beta_j - \log K.$$

Vanneste also shows that for an accurate estimate of L , we now only need $K \gg N$ and not, as in the brute-force case, $K \gg e^{cN}$ with $c > 0$.

3.5 Replica Trick

Just as the resampled Monte-Carlo method, the Replica Trick can be used to calculate the generalized Lyapunov exponent. Its attraction stems from the fact that the mean is taken directly of a random matrix, not of a nonlinear function. Hence the method is well suited for linear perturbations. The price to pay is that not all exponents $L(\mathcal{A}, \alpha)$, $\alpha \in \mathbb{R}$ can be calculated, but only $\alpha \in 2\mathbb{N}_+$.

Definition 5.

Let $\mathbf{A} \in \mathbb{R}^{d \times e}$. The vectorization of $\mathbf{A}$ by stacking its columns is defined as

$$\text{vec}(\mathbf{A})_i := \mathbf{A}_{(i-1 \operatorname{div} d)+1, (i-1 \operatorname{mod} e)+1}$$

for $1 \leq i \leq de$.

Definition 6.

Let $\mathbf{Z}, \mathbf{Y} \in \mathbb{R}^{d \times e}$, $\mathbf{W} \in \mathbb{R}^{f \times g}$. Then the Kronecker product $\mathbf{Z} \otimes \mathbf{W}$ of $\mathbf{Z}$ and $\mathbf{W}$ is defined as

$$(\mathbf{Z} \otimes \mathbf{W})_{i,j} := \mathbf{Z}_{(i-1 \operatorname{div} d)+1, (j-1 \operatorname{div} e)+1} \mathbf{W}_{(i-1 \operatorname{mod} f)+1, (j-1 \operatorname{mod} g)+1}$$

for $1 \leq i \leq df$ and $1 \leq j \leq eg$. We also define the k -th Kronecker power $\mathbf{Z}^{\otimes k}$ of $\mathbf{Z}$ as

$$\mathbf{Z}^{\otimes k} := \underbrace{\mathbf{Z} \otimes \mathbf{Z} \otimes \cdots \otimes \mathbf{Z}}_{k \text{ times}}.$$

Lemma 7.

Let $\mathbf{A}, \mathbf{B}, \mathbf{C} \in \mathbb{R}^{d \times d}$ and $\mathbf{I}$ the identity in $\mathbb{R}^d$. We have the following identities:

- $\operatorname{tr}(\mathbf{A} \otimes \mathbf{B}) = \operatorname{tr}(\mathbf{A}) \operatorname{tr}(\mathbf{B})$,
- $\operatorname{tr}(\mathbf{A}) = \operatorname{vec}(\mathbf{I}) \cdot \operatorname{vec}(\mathbf{A})$,
- $(\mathbf{A} \otimes \mathbf{C}) \operatorname{vec}(\mathbf{B}) = \operatorname{vec}(\mathbf{ABC}^T)$.

Lemma 8 (Replica trick).

Let $\mathcal{A}$ be a random square matrix. Then

$$L(\mathcal{A}, 2k) = \log \rho(\mathbb{E}[\mathcal{A}^{\otimes 2k}])$$

for $k \in \mathbb{N}$.

Proof. Let $\mathbf{A}$ a nonrandom square matrix. We write using the above identities

$$\begin{aligned} \|\mathbf{A}\|_F^{2k} &= \operatorname{tr}(\mathbf{AA}^T)^k \\ &= \operatorname{tr}((\mathbf{AA}^T)^{\otimes k}) \\ &= \operatorname{vec}(\mathbf{I}^{\otimes k}) \cdot \operatorname{vec}((\mathbf{AA}^T)^{\otimes k}) \\ &= \operatorname{vec}(\mathbf{I}^{\otimes k}) \cdot \operatorname{vec}(\mathbf{A}^{\otimes k} \mathbf{I}^{\otimes k} (\mathbf{A}^T)^{\otimes k}) \end{aligned}$$

$$= \text{vec}(\mathbf{I}^{\otimes k}) \cdot \mathbf{A}^{\otimes 2k} \text{vec}(\mathbf{I}^{\otimes k}).$$

Hence, if we call $\lambda_{j,\otimes 2k}$ and $v_{j,\otimes 2k}$ the eigenvalues and eigenvectors of $\mathbb{E}[\mathcal{A}^{\otimes 2k}]$, sorted in descending order with respect to their absolute value, we have

$$\begin{aligned} \mathbb{E} \left[\|\mathcal{A}^N\|_F^{2k} \right] &= \text{vec}(\mathbf{I}^{\otimes k}) \cdot \mathbb{E} \left[(\mathcal{A}^{\otimes 2k})^N \right] \text{vec}(\mathbf{I}^{\otimes k}) \\ &= \text{vec}(\mathbf{I}^{\otimes k}) \cdot \mathbb{E} [\mathcal{A}^{\otimes 2k}]^N \text{vec}(\mathbf{I}^{\otimes k}) \\ &= [\text{vec}(\mathbf{I}^{\otimes k}) \cdot v_{1,\otimes 2k}]^2 \lambda_{1,\otimes 2k}^N + \mathcal{O}(\lambda_{2,\otimes 2k}^N). \end{aligned}$$

Here, we assumed without loss of generality that $|\lambda_{1,\otimes 2k}| > |\lambda_{2,\otimes 2k}|$. Otherwise, we would only have a different constant multiplying the leading order term. Therefore, after taking the logarithm, dividing by N and letting $N \rightarrow \infty$, we obtain

$$L(\mathcal{A}, 2k) = \log \lambda_{1,\otimes 2k}.$$

Moreover, we also see that $\left(\frac{\lambda_{2,\otimes 2k}}{\lambda_{1,\otimes 2k}} \right)^N$ can be used to estimate the finite N variations of the generalized Lyapunov exponents. $\square$

The determination of the Lyapunov spectral radius is in general a hard problem. In what follows, we will limit ourselves to finding approximations numerically, as presented in Section 3.2, or estimates using the generalized Lyapunov exponent and the Replica Trick. While we have shown that the Lyapunov spectral radius can be bounded by

$$\varrho(\mathcal{A}) \leq \varrho_2(\mathcal{A}) \leq \sqrt{\rho(\mathbb{E}[\mathcal{A}^{\otimes 2}])},$$

the sharpness of the estimate is not always known and depends on the random matrix $\mathcal{A}$, as shown in Example 3. Therefore, in what follows, we complement analytic convergence estimates for the Lyapunov spectral radius with Monte-Carlo simulations to determine if the predicted behavior matches the observed convergence rates of fault-prone methods.

Chapter 4

Application of the Replica Trick to a Simple Iteration

In this section, we will turn to the application of the Replica Trick to the analysis of a fault-prone linear iterative method. Before we discuss the more complicated cases of the Fault-Prone Two Grid and Multigrid Methods, we illustrate the approach by limiting ourselves to the case of a fault-prone smoothener.

The iteration matrix for the fault-prone smoothener eq. (2.6) is given by

$$\mathcal{E}_\ell^S = I - \mathcal{X}_\ell^{(S)} N_\ell A_\ell,$$

Here, we take $\mathcal{X}_\ell^{(S)} = \text{diag}(\chi_1, \dots, \chi_{n_\ell})$ to be a random diagonal matrix of componentwise faults with $\mathbb{E}(\chi_j) = 1 - \varepsilon$ and $\text{Cov}(\chi_j, \chi_j) = \varepsilon(1 - \varepsilon)\delta_{ij}$. The behavior of the fault-prone smoothener is governed by the Lyapunov spectral radius $\varrho(\mathcal{E}_\ell^S)$. Our goal is to determine under which conditions the smoothener will stay convergent, i.e. $\varrho(\mathcal{E}_\ell^S) < 1$.

The following Lemma permits us to expand the second moment of a fault-prone iteration matrix in terms of expectations and variances of the fault matrices:

Lemma 9.

Let $\mathbf{B}_i \in \mathbb{R}^{n_i \times n_{i+1}}$, $i = 1, \dots, k$, be independent random matrices with $n_1 = n_{k+1}$. Then

$$\mathbb{E}[(\mathbf{I} - \mathbf{B}_1 \cdots \mathbf{B}_k)^{\otimes 2}] = \mathbb{E}[\mathbf{I} - \mathbf{B}_1 \cdots \mathbf{B}_k]^{\otimes 2} + \sum_{\mathbf{E} \in S} \mathbf{E}_1 \cdots \mathbf{E}_k,$$

where

$$S = \left[\bigotimes_{j=1}^k \{ \mathbb{E}[\mathbf{B}_j]^{\otimes 2}, \mathbb{V}[\mathbf{B}_j] \} \right] \setminus \{ (\mathbb{E}[\mathbf{B}_1]^{\otimes 2}, \dots, \mathbb{E}[\mathbf{B}_k]^{\otimes 2}) \}.$$

Proof. We multiply out to find

$$\begin{aligned} \mathbb{E}[(\mathbf{I} - \mathbf{B}_1 \cdots \mathbf{B}_k)^{\otimes 2}] &= \mathbb{E}[\mathbf{I}^{\otimes 2} - \mathbf{I} \otimes (\mathbf{B}_1 \cdots \mathbf{B}_k) - (\mathbf{B}_1 \cdots \mathbf{B}_k) \otimes \mathbf{I} + (\mathbf{B}_1 \cdots \mathbf{B}_k)^{\otimes 2}] \\ &= \mathbf{I}^{\otimes 2} - \mathbf{I} \otimes \mathbb{E}[\mathbf{B}_1 \cdots \mathbf{B}_k] - \mathbb{E}[\mathbf{B}_1 \cdots \mathbf{B}_k] \otimes \mathbf{I} + \mathbb{E}[(\mathbf{B}_1 \cdots \mathbf{B}_k)^{\otimes 2}] \\ &= \mathbb{E}[\mathbf{I} - \mathbf{B}_1 \cdots \mathbf{B}_k]^{\otimes 2} - \mathbb{E}[\mathbf{B}_1]^{\otimes 2} \cdots \mathbb{E}[\mathbf{B}_k]^{\otimes 2} + \mathbb{E}[\mathbf{B}_1^{\otimes 2}] \cdots \mathbb{E}[\mathbf{B}_k^{\otimes 2}]. \end{aligned}$$

In the last step, we used the independence of the random matrices. Now, because

$$\mathbb{E}[\mathbf{B}_j^{\otimes 2}] = \mathbb{V}[\mathbf{B}_j] + \mathbb{E}[\mathbf{B}_j]^{\otimes 2},$$

we can combine the last two terms in the expression above to obtain the desired result. $\square$

Theorem 10.

Let $\mathbf{E}_\ell^S = \mathbf{I} - \mathbf{N}_\ell \mathbf{A}_\ell$ be the iteration matrix for any convergent smootheners. Then, the corresponding fault-prone smootheners satisfies

$$\varrho(\mathcal{E}_\ell^S) \leq \sqrt{(1 - \varepsilon) \|\mathbf{E}_\ell^S\|_2^2 + \varepsilon (\|\mathbf{E}_\ell^S\|_2 + \|\mathbf{N}_\ell \mathbf{A}_\ell\|_2)^2},$$

for all $\varepsilon \in [0, 1]$.

Proof. Let $\mathbf{B} = \mathcal{X}^{(S)} \mathbf{N} \mathbf{A}$ and dispense with subscripts and superscripts for the duration of the proof. Then

$$\begin{aligned} \mathbb{E}[\mathcal{E}^{\otimes 2}] &= \mathbb{E}[\mathbf{I}^{\otimes 2} - \mathbf{B} \otimes \mathbf{I} - \mathbf{I} \otimes \mathbf{B} + \mathbf{B}^{\otimes 2}] \\ &= \mathbf{I}^{\otimes 2} - \mathbb{E}[\mathbf{B}] \otimes \mathbf{I} - \mathbf{I} \otimes \mathbb{E}[\mathbf{B}] + \mathbb{E}[\mathbf{B}]^{\otimes 2} \end{aligned}$$

$$\begin{aligned}
& -\mathbb{E}[\mathcal{B}]^{\otimes 2} + \mathbb{E}[\mathcal{B}^{\otimes 2}] \\
& = \mathbb{E}[\mathcal{E}]^{\otimes 2} + \mathbb{V}[\mathcal{B}]
\end{aligned} \tag{4.1}$$

and

$$\mathbb{V}[\mathcal{B}] = \mathbb{V}\left[\mathcal{X}^{(S)}\right](\mathbf{N}\mathbf{A})^{\otimes 2} = \varepsilon(1 - \varepsilon)\mathbf{K}(\mathbf{N}\mathbf{A})^{\otimes 2},$$

where

$$\mathbf{K} = \text{blockdiag}\left(\vec{e}^{(j)} \otimes (\vec{e}^{(j)})^T, j = 1, \dots, n\right)$$

and $\vec{e}^{(j)}$ the j -th canonical unit vector in $\mathbb{R}^n$. We will derive a more general form of eq. (4.1) in Lemma 9 in the appendix. Since $\|\mathbf{C}^{\otimes 2}\|_2 = \|\mathbf{C}\|_2^2$ and $\|\mathbf{K}_\ell \mathbf{C}\|_2 \leq \|\mathbf{C}\|_2$ for any compatible matrix $\mathbf{C}$, we obtain

$$\|\mathbb{V}[\mathcal{B}]\|_2 = \varepsilon(1 - \varepsilon) \|\mathbf{K}(\mathbf{N}\mathbf{A})^{\otimes 2}\|_2 \leq \varepsilon(1 - \varepsilon) \|\mathbf{N}\mathbf{A}\|_2^2$$

and

$$\|\mathbb{E}[\mathcal{E}]^{\otimes 2}\|_2 = \|\mathbb{E}[\mathcal{E}]\|_2^2 = \|\mathbf{I} - (1 - \varepsilon)\mathbf{N}\mathbf{A}\|_2^2.$$

Consequently,

$$\begin{aligned}
\|\mathbb{E}[\mathcal{E}^{\otimes 2}]\|_2 & \leq \|\mathbf{I} - (1 - \varepsilon)\mathbf{N}\mathbf{A}\|_2^2 + \varepsilon(1 - \varepsilon) \|\mathbf{N}\mathbf{A}\|_2^2 \\
& = \|\mathbf{E} + \varepsilon\mathbf{N}\mathbf{A}\|_2^2 + \varepsilon(1 - \varepsilon) \|\mathbf{N}\mathbf{A}\|_2^2 \\
& \leq (\|\mathbf{E}\|_2 + \varepsilon \|\mathbf{N}\mathbf{A}\|_2)^2 + \varepsilon(1 - \varepsilon) \|\mathbf{N}\mathbf{A}\|_2^2 \\
& = (1 - \varepsilon) \|\mathbf{E}\|_2^2 + \varepsilon(\|\mathbf{E}\|_2 + \|\mathbf{N}\mathbf{A}\|_2)^2
\end{aligned}$$

and the result follows thanks to Lemma 8. $\square$

In order to illustrate Theorem 10, we consider the situation where the smoothener is taken to be a convergent relaxation scheme so that $\|\mathbf{E}_\ell^S\|_2 = \Gamma \in [0, 1)$. The triangle inequality gives $\|\mathbf{N}_\ell \mathbf{A}_\ell\|_2 \leq 1 + \Gamma$ and hence, thanks to Theorem 10, we deduce that

$$\varrho(\mathcal{E}_\ell^S) \leq \sqrt{\Gamma^2 + \varepsilon(1 + 4\Gamma + 3\Gamma^2)} = \Gamma + \mathcal{O}(\varepsilon).$$

As a consequence, we find that a sufficient condition for the fault-prone smoothener to remain convergent is that the fault rate is sufficiently small: $\varepsilon < (1 - \Gamma)/(1 + 3\Gamma)$.

Chapter 5

Analysis of the Fault-Prone Two Grid Method

We now turn to the analysis of the Fault-Prone Two Grid Method as given in Algorithm 3. We use subscripts C for coarse-level quantities, and F for fine-level quantities.

The convergence analysis will be carried out with respect to the energy norm:

Definition 11 (Energy norms).

For matrices $\mathbf{Z} \in \mathbb{R}^{n_\ell \times n_\ell}$, we define the usual matrix energy norm $\|\mathbf{Z}\|_A$ as well as the double energy norm $\|\mathbf{Z}\|_{A^2}$ to be $\|\mathbf{Z}\|_A = \left\| \mathbf{A}_\ell^{\frac{1}{2}} \mathbf{Z} \mathbf{A}_\ell^{-\frac{1}{2}} \right\|_2$, and $\|\mathbf{Z}\|_{A^2} = \|\mathbf{A}_\ell \mathbf{Z} \mathbf{A}_\ell^{-1}\|_2$. For matrices $\mathbf{W} \in \mathbb{R}^{n_\ell^2 \times n_\ell^2}$, we define the tensor energy norm $\|\mathbf{W}\|_A$ and the tensor double energy norm $\|\mathbf{W}\|_{A^2}$ to be $\|\mathbf{W}\|_A = \left\| \left(\mathbf{A}_\ell^{\frac{1}{2}} \right)^{\otimes 2} \mathbf{W} \left(\mathbf{A}_\ell^{-\frac{1}{2}} \right)^{\otimes 2} \right\|_2$, and $\|\mathbf{W}\|_{A^2} = \left\| \mathbf{A}_\ell^{\otimes 2} \mathbf{W} \left(\mathbf{A}_\ell^{-1} \right)^{\otimes 2} \right\|_2$. In all cases $\|\bullet\|_2$ is the spectral norm.

The first main result of the present work concerns the convergence of the Fault-Prone Two-Grid Algorithm for the special case of all $\mathcal{X}_{\bullet}^{(\bullet)}$ having independent Bernoulli distributed diagonal entries with probability of fault ε :

Theorem 12.

Let $\mathcal{E}^{TG} = (\mathcal{E}^{S,post})^{\nu_2} \mathcal{E}^{CG} (\mathcal{E}^{S,pre})^{\nu_1}$ be the iteration matrix of the Fault-Prone Two Grid

Method (Algorithm 3) with Jacobi smoothener and component-wise faults of rate ε in prolongation, restriction, residual and in the smoothener, and let $\mathbf{E}^{TG}$ be its fault-free equivalent. Assume that $\Omega \in C^2$ or Ω a convex polyhedron, $\mathcal{T}_C$ quasiuniform and that Assumptions 1 to 4 hold. Then

$$\varrho(\mathcal{E}^{TG}) \leq \|\mathbf{E}^{TG}\|_A + C \begin{cases} \varepsilon n_F^{(4-d)/2d} & d < 4, \\ \varepsilon \sqrt{\log n_F} & d = 4, \\ \varepsilon & d > 4, \end{cases} \quad (5.1)$$

where $\|\bullet\|_A$ is the matrix energy norm and the constant C is independent of ε and n_F .

Before we turn to the proof of Theorem 12, we record several useful definition and results.

Definition 13.

Let $\mathbf{Z}, \mathbf{Y} \in \mathbb{R}^{n \times m}$, $n, m \in \mathbb{N}$. The element wise or Hadamard product $\mathbf{Z} \circ \mathbf{Y} \in \mathbb{R}^{n \times m}$ of $\mathbf{Z}$ and $\mathbf{Y}$ is defined as

$$(\mathbf{Z} \circ \mathbf{Y})_{ij} = \mathbf{Z}_{ij} \mathbf{Y}_{ij}$$

for $1 \leq i \leq n$, $1 \leq j \leq m$, and the k -th Hadamard power $\mathbf{Z}^{\circ k} \in \mathbb{R}^{n \times m}$ of $\mathbf{Z}$ as

$$\mathbf{Z}^{\circ k} := \underbrace{\mathbf{Z} \circ \mathbf{Z} \circ \dots \circ \mathbf{Z}}_{k \text{ times}}.$$

Lemma 14 (Horn and Johnson [66]).

Let $\mathbf{H} = \begin{pmatrix} \mathbf{A} & \mathbf{B} \\ \mathbf{B}^T & \mathbf{C} \end{pmatrix} \in \mathbb{R}^{(p+q) \times (p+q)}$ be symmetric, $\mathbf{A} \in \mathbb{R}^{p \times p}$, $\mathbf{B} \in \mathbb{R}^{p \times q}$, $\mathbf{C} \in \mathbb{R}^{q \times q}$. Then $\mathbf{H}$ is positive (semi-) definite if and only if $\mathbf{A}$ and $\mathbf{C}$ are positive (semi-) definite and there exists a contraction $\mathbf{X} \in \mathbb{R}^{p \times q}$ (i.e. all singular values are smaller or equal than 1) such that $\mathbf{B} = \mathbf{A}^{\frac{1}{2}} \mathbf{X} \mathbf{C}^{\frac{1}{2}}$.

Lemma 15 (Horn and Johnson [67]).

Let $\mathbf{Z}, \mathbf{Y}$ be symmetric positive (semi-) definite matrices. Then $\mathbf{Z} \circ \mathbf{Y}$ is symmetric positive (semi-) definite.

We now prove the following estimates for the Hadamard power of matrices:

Lemma 16.

Let $\mathbf{Z} \in \mathbb{R}^{n \times m}$. Then

$$\|\mathbf{Z}^{\circ 2}\|_2 \leq \max_{i=1,\dots,m} \|\mathbf{Z}\vec{e}_m^{(i)}\|_2 \max_{j=1,\dots,n} \|\mathbf{Z}^T \vec{e}_n^{(j)}\|_2 \leq \|\mathbf{Z}\|_2^2,$$

where $\{\vec{e}_m^{(i)}\}$ and $\{\vec{e}_n^{(j)}\}$ are the canonical unit basis vectors in $\mathbb{R}^m$ and $\mathbb{R}^n$ respectively.

Proof. Both the matrix $\begin{pmatrix} \mathbf{Z} & \mathbf{I}_n \end{pmatrix}^T \begin{pmatrix} \mathbf{Z} & \mathbf{I}_n \end{pmatrix}$ and $\begin{pmatrix} \mathbf{I}_m & \mathbf{Z}^T \end{pmatrix}^T \begin{pmatrix} \mathbf{I}_m & \mathbf{Z}^T \end{pmatrix}$ are symmetric positive semi-definite. Using Lemma 15, their Hadamard product

$$\begin{pmatrix} (\mathbf{Z}^T \mathbf{Z}) \circ \mathbf{I}_m & (\mathbf{Z}^T)^{\circ 2} \\ \mathbf{Z}^{\circ 2} & \mathbf{I}_n \circ (\mathbf{Z} \mathbf{Z}^T) \end{pmatrix}$$

is also positive semi-definite. Hence, by Lemma 14,

$$\|\mathbf{Z}^{\circ 2}\|_2^2 \leq \|(\mathbf{Z} \mathbf{Z}^T) \circ \mathbf{I}_n\|_2 \|(\mathbf{Z}^T \mathbf{Z}) \circ \mathbf{I}_m\|_2 = \max_i \|\mathbf{Z}\vec{e}_m^{(i)}\|_2^2 \max_i \|\mathbf{Z}^T \vec{e}_n^{(j)}\|_2^2.$$

The second inequality follows trivially. $\square$

Proof of Theorem 12mg. We set

$$\mathbb{E}[\mathcal{X}^{(\bullet)}] = e^{(\bullet)} \mathbf{I}, \quad \mathbb{V}[\mathcal{X}^{(\bullet)}] = \mathbf{V}^{(\bullet)}.$$

Using Lemma 9 the second moment of the fault-prone coarse grid correction and smoothener can be written as

$$\begin{aligned} \mathbb{E}[(\mathcal{E}^{CG})^{\otimes 2}] &= \mathbb{E}[\mathcal{E}^{CG}]^{\otimes 2} + \mathbf{C}^{(P)} + \mathbf{C}^{(R)} + \mathbf{C}^{(\rho)} \\ &\quad + \mathbf{C}^{(P,R)} + \mathbf{C}^{(P,\rho)} + \mathbf{C}^{(R,\rho)} + \mathbf{C}^{(P,R,\rho)}, \end{aligned} \tag{5.2}$$

$$\mathbb{E}[(\mathcal{E}^S)^{\otimes 2}] = \mathbb{E}[\mathcal{E}^S]^{\otimes 2} + \mathbf{C}^{(S)}, \tag{5.3}$$

with

$$\mathbb{E}[\mathcal{E}^{CG}] = \mathbf{E}^{CG} + (1 - e^{(P)} e^{(R)} e^{(\rho)}) (\mathbf{I} - \mathbf{E}^{CG}),$$

$$\begin{aligned}
\mathbf{C}^{(P)} &= (e^{(R)} e^{(\rho)})^2 \mathbf{V}^{(P)} (\mathbf{P} \mathbf{A}_C^{-1} \mathbf{R} \mathbf{A}_F)^{\otimes 2}, \\
\mathbf{C}^{(R)} &= (e^{(P)} e^{(\rho)})^2 (\mathbf{P} \mathbf{A}_C^{-1})^{\otimes 2} \mathbf{V}^{(R)} (\mathbf{R} \mathbf{A}_F)^{\otimes 2}, \\
\mathbf{C}^{(\rho)} &= (e^{(P)} e^{(R)})^2 (\mathbf{P} \mathbf{A}_C^{-1} \mathbf{R})^{\otimes 2} \mathbf{V}^{(\rho)} \mathbf{A}_F^{\otimes 2}, \\
\mathbf{C}^{(P,R)} &= (e^{(\rho)})^2 \mathbf{V}^{(P)} (\mathbf{P} \mathbf{A}_C^{-1})^{\otimes 2} \mathbf{V}^{(R)} (\mathbf{R} \mathbf{A}_F^{\otimes 2})^{\otimes 2}, \\
\mathbf{C}^{(P,\rho)} &= (e^{(R)})^2 \mathbf{V}^{(P)} (\mathbf{P} \mathbf{A}_C^{-1} \mathbf{R})^{\otimes 2} \mathbf{V}^{(\rho)} \mathbf{A}_F^{\otimes 2}, \\
\mathbf{C}^{(R,\rho)} &= (e^{(P)})^2 (\mathbf{P} \mathbf{A}_C^{-1})^{\otimes 2} \mathbf{V}^{(R)} \mathbf{R}^{\otimes 2} \mathbf{V}^{(\rho)} \mathbf{A}_F^{\otimes 2}, \\
\mathbf{C}^{(P,R,\rho)} &= \mathbf{V}^{(P)} (\mathbf{P} \mathbf{A}_C^{-1})^{\otimes 2} \mathbf{V}^{(R)} \mathbf{R}^{\otimes 2} \mathbf{V}^{(\rho)} \mathbf{A}_F^{\otimes 2}, \\
\mathbb{E} [\boldsymbol{\mathcal{E}}^S] &= \mathbf{E}^S + e^{(S)} (\mathbf{I} - \mathbf{E}^S), \\
\mathbf{C}^{(S)} &= \mathbf{V}^{(S)} (\mathbf{N} \mathbf{A}_F)^{\otimes 2}.
\end{aligned}$$

We have that

$$e^{(\bullet)} = 1 - \varepsilon, \quad \mathbf{V}^{(P)} = \mathbf{V}^{(\rho)} = \mathbf{V}^{(S)} = \varepsilon(1 - \varepsilon) \mathbf{K}_F, \quad \mathbf{V}^{(R)} = \varepsilon(1 - \varepsilon) \mathbf{K}_C,$$

with

$$\begin{aligned}
\mathbf{K} &= \text{blockdiag} \left(\vec{e}_F^{(i)} \otimes (\vec{e}_F^{(i)})^T, i = 1, \dots, n_F \right), \\
\mathbf{K}_C &= \text{blockdiag} \left(\vec{e}_C^{(i)} \otimes (\vec{e}_C^{(i)})^T, i = 1, \dots, n_C \right).
\end{aligned}$$

Here, $\vec{e}_F^{(i)}$ and $\vec{e}_C^{(i)}$ are the canonical unit basis vectors of $\mathbb{R}^{n_F}$ and $\mathbb{R}^{n_C}$ respectively.

Adding and subtracting $\mathbb{E} [\boldsymbol{\mathcal{E}}^{TG}]^{\otimes 2}$ from $\mathbb{E} [(\boldsymbol{\mathcal{E}}^{TG})^{\otimes 2}]$, we estimate using the energy norm on the tensor space

$$\begin{aligned}
\rho \left(\mathbb{E} [(\boldsymbol{\mathcal{E}}^{TG})^{\otimes 2}] \right) &\leq \left\| \mathbb{E} [(\boldsymbol{\mathcal{E}}^{TG})^{\otimes 2}] \right\|_A \\
&\leq \left\| \mathbb{E} [\boldsymbol{\mathcal{E}}^{TG}] \right\|_A^2 + \left\| \mathbb{E} [((\boldsymbol{\mathcal{E}}^{S,\text{post}})^{\nu_2} \boldsymbol{\mathcal{E}}^{CG} (\boldsymbol{\mathcal{E}}^{S,\text{pre}})^{\nu_1})^{\otimes 2}] \right\|_A \\
&\quad - \left\| \mathbb{E} [(\boldsymbol{\mathcal{E}}^{S,\text{post}})^{\nu_2} \boldsymbol{\mathcal{E}}^{CG} (\boldsymbol{\mathcal{E}}^{S,\text{pre}})^{\nu_1}] \right\|_A^2.
\end{aligned}$$

We then use the subadditivity of $\|\bullet\|_A$ and equations eq. (5.2) and eq. (5.3) to write

$$\begin{aligned}
\rho \left(\mathbb{E} [(\boldsymbol{\mathcal{E}}^{TG})^{\otimes 2}] \right) &\leq \left\| \mathbb{E} [\boldsymbol{\mathcal{E}}^{TG}] \right\|_A^2 \\
&\quad + \left(\left\| \mathbb{E} [\boldsymbol{\mathcal{E}}^{CG}] \right\|_A^2 + \left\| \mathbf{C}^{(P)} \right\|_A + \left\| \mathbf{C}^{(R)} \right\|_A + \left\| \mathbf{C}^{(\rho)} \right\|_A \right. \\
&\quad \left. + \left\| \mathbf{C}^{(P,R)} \right\|_A + \left\| \mathbf{C}^{(P,\rho)} \right\|_A + \left\| \mathbf{C}^{(R,\rho)} \right\|_A + \left\| \mathbf{C}^{(P,R,\rho)} \right\|_A \right)
\end{aligned}$$

$$\begin{aligned}
& \times \left(\|\mathbb{E}[\boldsymbol{\mathcal{E}}^S]\|_A^2 + \|\boldsymbol{C}^{(S)}\|_A \right)^{\nu_1 + \nu_2} \\
& - \|\mathbb{E}[\boldsymbol{\mathcal{E}}^{CG}]\|_A^2 \|\mathbb{E}[\boldsymbol{\mathcal{E}}^S]\|_A^{2(\nu_1 + \nu_2)}.
\end{aligned} \tag{5.4}$$

First, we estimate the terms involving only first moments of the fault matrices. Using that $\boldsymbol{E}^{CG}$ is an A -orthogonal projection and that the damped Jacobi smootheners are convergent, we find

$$\begin{aligned}
\|\mathbb{E}[\boldsymbol{\mathcal{E}}^{CG}]\|_A & \leq \|\boldsymbol{E}^{CG}\|_A + (1 - (1 - \varepsilon)^3) \|\boldsymbol{I} - \boldsymbol{E}^{CG}\|_A \leq 1 + C\varepsilon, \\
\|\mathbb{E}[\boldsymbol{\mathcal{E}}^S]\|_A & \leq \|\boldsymbol{E}^S\|_A + \varepsilon \|\boldsymbol{I} - \boldsymbol{E}^S\|_A \leq 1 + C\varepsilon,
\end{aligned}$$

and therefore

$$\begin{aligned}
\|\mathbb{E}[\boldsymbol{\mathcal{E}}^{TG}]\|_A & \leq \|\boldsymbol{E}^{TG}\|_A + \|\mathbb{E}[\boldsymbol{\mathcal{E}}^S]\|_A^{\nu_1 + \nu_2} \|\mathbb{E}[\boldsymbol{\mathcal{E}}^{CG}]\|_A - \|\boldsymbol{E}^S\|_A^{\nu_1 + \nu_2} \|\boldsymbol{E}^{CG}\|_A \\
& = \|\boldsymbol{E}^{TG}\|_A + C\varepsilon.
\end{aligned}$$

Next, we estimate all the terms $\boldsymbol{C}^{(\bullet)}$ involving variances of the fault matrices. We bound

$$\begin{aligned}
\frac{1}{\varepsilon} \|\boldsymbol{C}^{(\rho)}\|_A & \leq \rho \left[\left(\boldsymbol{A}_F^{\frac{1}{2}} \boldsymbol{P} \boldsymbol{A}_C^{-1} \boldsymbol{R} \right)^{\otimes 2} \boldsymbol{K}_F \boldsymbol{A}_F^{\otimes 2} \boldsymbol{K}_F \left(\boldsymbol{P} \boldsymbol{A}_C^{-1} \boldsymbol{R} \boldsymbol{A}_F^{\frac{1}{2}} \right)^{\otimes 2} \right]^{\frac{1}{2}} \\
& = \rho \left[\left(\boldsymbol{P} \boldsymbol{A}_C^{-1} \boldsymbol{R} \boldsymbol{A}_F \boldsymbol{P} \boldsymbol{A}_C^{-1} \boldsymbol{R} \right)^{\otimes 2} \boldsymbol{K}_F \boldsymbol{A}_F^{\otimes 2} \boldsymbol{K}_F \right]^{\frac{1}{2}} \\
& = \rho \left[\left(\boldsymbol{P} \boldsymbol{A}_C^{-1} \boldsymbol{R} \right)^{\circ 2} \boldsymbol{A}_F^{\circ 2} \right]^{\frac{1}{2}} \\
& \leq \left\| \left(\boldsymbol{P} \boldsymbol{A}_C^{-1} \boldsymbol{R} \right)^{\circ 2} \right\|_2^{\frac{1}{2}} \left\| \boldsymbol{A}_F^{\circ 2} \right\|_2^{\frac{1}{2}}.
\end{aligned}$$

Here, we used that $\boldsymbol{K}_F = \boldsymbol{K}_F^2$ and that for compatible $\boldsymbol{Z}$

$$\begin{aligned}
(\boldsymbol{K}_F \boldsymbol{Z}^{\otimes 2} \boldsymbol{K}_F)_{in+p, jn+q} & = (\boldsymbol{K}_F)_{in+p, in+p} \boldsymbol{Z}_{ij} \boldsymbol{Z}_{pq} (\boldsymbol{K}_F)_{jn+q, jn+q} \\
& = \begin{cases} (\boldsymbol{Z}^{\circ 2})_{ij} & \text{if } i = p, j = q, \\ 0 & \text{else} \end{cases}
\end{aligned}$$

for $1 \leq i, j, p, q \leq n$. Similarly, we find

$$\frac{1}{\varepsilon} \|\boldsymbol{C}^{(P)}\|_A \leq \left\| \boldsymbol{A}_F^{\circ 2} \right\|_2^{\frac{1}{2}} \left\| \left(\boldsymbol{P} \boldsymbol{A}_C^{-1} \boldsymbol{R} \right)^{\circ 2} \right\|_2^{\frac{1}{2}},$$

$$\begin{aligned}
\frac{1}{\varepsilon} \left\| \mathbf{C}^{(R)} \right\|_A &\leq \left\| \mathbf{A}_C^{\circ 2} \right\|_2^{\frac{1}{2}} \left\| (\mathbf{A}_C^{-1})^{\circ 2} \right\|_2^{\frac{1}{2}}, \\
\frac{1}{\varepsilon^2} \left\| \mathbf{C}^{(P,R)} \right\|_A &\leq \left\| \mathbf{A}_F^{\circ 2} \right\|_2^{\frac{1}{2}} \left\| \mathbf{A}_C^{\circ 2} \right\|_2^{\frac{1}{2}} \left\| (\mathbf{A}_C^{-1} \mathbf{R})^{\circ 2} \right\|_2, \\
\frac{1}{\varepsilon^2} \left\| \mathbf{C}^{(P,\rho)} \right\|_A &\leq \left\| \mathbf{A}_F^{\circ 2} \right\|_2 \left\| (\mathbf{P} \mathbf{A}_C^{-1} \mathbf{R})^{\circ 2} \right\|_2, \\
\frac{1}{\varepsilon^2} \left\| \mathbf{C}^{(R,\rho)} \right\|_A &\leq \left\| \mathbf{P}^{\circ 2} \right\|_2 \left\| \mathbf{R}^{\circ 2} \right\|_2 \left\| \mathbf{A}_F^{\circ 2} \right\|_2^{\frac{1}{2}} \left\| (\mathbf{A}_C^{-1})^{\circ 2} \right\|_2^{\frac{1}{2}}, \\
\frac{1}{\varepsilon^3} \left\| \mathbf{C}^{(P,R,\rho)} \right\|_A &\leq \left\| \mathbf{P}^{\circ 2} \right\|_2^2 \left\| \mathbf{A}_F^{\circ 2} \right\|_2 \left\| (\mathbf{A}_C^{-1} \mathbf{R})^{\circ 2} \right\|_2, \\
\frac{1}{\varepsilon} \left\| \mathbf{C}^{(S)} \right\|_A &\leq \left\| (\mathbf{N} \mathbf{A}_F)^{\circ 2} \right\|_2.
\end{aligned}$$

In the last inequality, we used that $\mathbf{V}^{(S)}$ and $\mathbf{N}^{\otimes 2}$ commute. Using Lemma 16 and Assumptions 2 and 4 we find

$$\begin{aligned}
\left\| \mathbf{P}^{\circ 2} \right\|_2 &\leq \left\| \mathbf{P} \right\|_2^2 \leq C, \\
\left\| \mathbf{A}_F^{\circ 2} \right\|_2 &\leq \left\| \mathbf{A}_F \right\|_2^2, \\
\left\| \mathbf{A}_C^{\circ 2} \right\|_2 &\leq \left\| \mathbf{A}_C \right\|_2^2 \leq \left\| \mathbf{A}_F \right\|_2^2, \\
\left\| (\mathbf{P} \mathbf{A}_C^{-1} \mathbf{R})^{\circ 2} \right\|_2 &\leq \max_i \left\| \mathbf{P} \mathbf{A}_C^{-1} \mathbf{R} \vec{e}_F^{(i)} \right\|_2^2 \\
&\leq \left\| \mathbf{P} \mathbf{A}_C^{-1} \mathbf{R} \mathbf{A}_F \right\|_2^2 \max_i \left\| \mathbf{A}_F^{-1} \vec{e}_F^{(i)} \right\|_2^2 \\
&\leq C \max_i \left\| \mathbf{A}_F^{-1} \vec{e}_F^{(i)} \right\|_2^2, \\
\left\| (\mathbf{A}_C^{-1} \mathbf{R})^{\circ 2} \right\|_2 &\leq \max_i \left\| \mathbf{P} \mathbf{A}_C^{-1} \vec{e}_C^{(i)} \right\|_2 \max_j \left\| \mathbf{A}_C^{-1} \mathbf{R} \vec{e}_F^{(j)} \right\|_2 \\
&\leq C \max_i \left\| \mathbf{A}_C^{-1} \vec{e}_C^{(i)} \right\|_2 \left\| \mathbf{A}_C^{-1} \mathbf{R} \mathbf{A}_F \right\|_2 \max_j \left\| \mathbf{A}_F^{-1} \vec{e}_F^{(j)} \right\|_2 \\
&\leq C \max_i \left\| \mathbf{A}_C^{-1} \vec{e}_C^{(i)} \right\|_2 \max_j \left\| \mathbf{A}_F^{-1} \vec{e}_F^{(j)} \right\|_2, \\
\left\| (\mathbf{N} \mathbf{A}_F)^{\circ 2} \right\|_2 &\leq \left\| \mathbf{N} \mathbf{A}_F \right\|_2^2 \leq C.
\end{aligned}$$

Now, because $\mathbf{A}_F$ is the finite element discretization of a second order PDE over a quasi-uniform mesh with mesh size h , we obtain (see Theorem 9.11 in [50]),

$$\left\| \mathbf{A}_F \right\|_2 \leq Ch^{d-2}.$$

Since

$$\left\| \mathbf{A}_F^{-1} \vec{e}_F^{(i)} \right\|_2 \leq C h^{-\frac{d}{2}} \|u_F\|_{L^2},$$

where $u_F \in V_F$ is given by

$$a(u_F, v) = v(\vec{x}_i), \quad \forall v \in V_F,$$

we can apply Lemma 17 to obtain bounds for $\max_j \left\| \mathbf{A}_F^{-1} \vec{e}_F^{(j)} \right\|_2$ (and equally for $\max_i \left\| \mathbf{A}_C^{-1} \vec{e}_C^{(i)} \right\|_2$).

$$\begin{aligned} \frac{1}{\varepsilon} \left\| \mathbf{C}^{(P)} \right\|_A, \frac{1}{\varepsilon} \left\| \mathbf{C}^{(R)} \right\|_A, \frac{1}{\varepsilon} \left\| \mathbf{C}^{(\rho)} \right\|_A, \frac{1}{\varepsilon^2} \left\| \mathbf{C}^{(R,\rho)} \right\|_A &\leq C \begin{cases} h^{\frac{d-4}{2}} & d < 4, \\ \sqrt{1 + |\log h|} & d = 4, \\ 1 & d > 4, \end{cases} \\ \frac{1}{\varepsilon^2} \left\| \mathbf{C}^{(P,R)} \right\|_A, \frac{1}{\varepsilon^2} \left\| \mathbf{C}^{(P,\rho)} \right\|_A, \frac{1}{\varepsilon^3} \left\| \mathbf{C}^{(P,R,\rho)} \right\|_A &\leq C \begin{cases} h^{d-4} & d < 4, \\ (1 + |\log h|) & d = 4, \\ 1 & d > 4, \end{cases} \\ \frac{1}{\varepsilon} \left\| \mathbf{C}^{(S)} \right\|_A &\leq C. \end{aligned}$$

Therefore, we obtain from eq. (5.4)

$$\rho \left(\mathbb{E} \left[(\boldsymbol{\mathcal{E}}^{TG})^{\otimes 2} \right] \right) \leq \left\| \mathbf{E}^{TG} \right\|_A^2 + C \begin{cases} \varepsilon^2 h^{d-4} & d < 4, \\ \varepsilon^2 (1 + |\log h|) & d = 4, \\ \varepsilon^2 & d > 4, \end{cases}$$

and hence

$$\varrho(\boldsymbol{\mathcal{E}}^{TG}) \leq \left\| \mathbf{E}^{TG} \right\|_A + C \begin{cases} \varepsilon n_F^{\frac{4-d}{2d}} & d < 4, \\ \varepsilon \sqrt{\log n_F} & d = 4, \\ \varepsilon & d > 4, \end{cases}$$

where we used that $n_F \approx h^{-d}$. □

In order to conclude, we need the following technical estimate:

Lemma 17.

Let $\Omega \in C^2$ or Ω a convex polyhedron and let $u_F \in V_F$ be the unique solution of

$$a(u_F, v) = v(\vec{x}_i), \quad \forall v \in V_F.$$

Then

$$\|u_F\|_{L^2} \leq C \begin{cases} 1 & d < 4, \\ (1 + |\log h|)^{\frac{1}{2}} & d = 4, \\ h^{2-\frac{d}{2}} & d > 4, \end{cases}$$

where C is a constant independent of h .

Proof. Let $f \in V_F$ be the unique function that corresponds to the load and write $f = \vec{F} \cdot \phi$ with $\vec{F}$ its coefficient vector and ϕ the vector of shape functions. Then

$$\vec{e}_F^{(i)} = (\phi_F, f)_{L^2} = (\phi_F, \phi_F)_{L^2} \cdot \vec{F} = \mathbf{M}_F \vec{F}.$$

Since by Theorem 9.8 in [50]

$$Ch^d \mathbf{I} \leq \mathbf{M}_F \leq Ch^d \mathbf{I}, \quad (5.5)$$

we have

$$\|f\|_{H^m} \leq Ch^{-d} \left\| \phi_F^{(i)} \right\|_{H^m}. \quad (5.6)$$

Now u_F is an approximation to $u \in V$ that solves

$$a(u, v) = (f, v)_{L^2}, \quad \forall v \in V.$$

Since $f \in V \subset L^2(\Omega)$, we find by the Aubin-Nitsche Lemma [50] that

$$\|u_F - u\|_{L^2} \leq Ch^2 \|f\|_{L^2}.$$

Consider the solution $u^* \in V$ to the dual problem

$$a(v, u^*) = (u, v)_{L^2}, \quad \forall v \in V.$$

By elliptic regularity [50], we have

$$\|u^*\|_{H^2} \leq C \|u\|_{L^2}.$$

Moreover,

$$(u, u)_{L^2} = a(u, u^*) = (f, u^*)_{L^2},$$

so

$$\|u\|_{L^2}^2 \leq \|f\|_{H^{-2}} \|u^*\|_{H^2} \leq C \|f\|_{H^{-2}} \|u\|_{L^2},$$

and hence $\|u\|_{L^2} \leq C \|f\|_{H^{-2}}$. Therefore, by triangle inequality and eq. (5.6), we find

$$\|u_F\|_{L^2} \leq Ch^2 \|f\|_{L^2} + C \|f\|_{H^{-2}} \leq Ch^{-d} \left(h^2 \left\| \phi_F^{(i)} \right\|_{L^2} + \left\| \phi_F^{(i)} \right\|_{H^{-2}} \right).$$

Applying the estimates for $\left\| \phi_F^{(i)} \right\|_{H^m}$ from Theorem 4.8 in [10], we obtain

$$\|u_F\|_{L^2} \leq C \begin{cases} 1 & d < 4, \\ (1 + |\log h|)^{\frac{1}{2}} & d = 4, \\ h^{2-\frac{d}{2}} & d > 4. \end{cases}$$

□

If the probability of a fault occurring is zero, i.e. $\varepsilon = 0$, then the bound eq. (5.1) reduces to the quantity $\|\mathbf{E}^{TG}\|_A$ which governs the convergence of the fault-free Two Grid Algorithm. Under the foregoing assumptions on the finite element discretization, it is known [20, 21, 63, 95] that $\|\mathbf{E}^{TG}\|_A \leq C_{TG} < 1$, where C_{TG} is a positive constant independent of n_F . If $\varepsilon \in (0, 1)$, then the bound eq. (5.1) depends

on the failure probability ε , the number of unknowns n_F and the spatial dimension d of the underlying discretization.

In the simplest case, in which the spatial dimension $d > 4$, the Fault-Prone Two Grid Method will converge uniformly in n_F provided that the failure probability ε is sufficiently small, i.e. $\varepsilon \in [0, (1 - C_{TG})/C) \subset [0, 1)$. The estimate in the cases of most practical interest, in which $d < 4$, is less encouraging with the bound depending on the size n_F of the problem being solved. In particular, the estimate means that no matter how small the probability ε of a failure may be, as the size of the problem being solved grows, the bound will exceed unity suggesting that the Two Grid Method will fail to converge *at all*.

5.1 Numerical Verification of Non-Resilience

In order to illustrate this behavior and to test the sharpness of the estimate eq. (5.1), we apply the Fault-Prone Two Grid Method to the system arising from a discretization of the Poisson problem on a square using piecewise linear functions on a uniform mesh of triangles (Figure 5.2).

We wish to investigate the behavior as the number of degrees of freedom n_F and the fault rate ε are varied. In order to do this, it will be necessary to simulate the faults that might be encountered on an actual system. To verify the scaling law eq. (5.1), we inject faults at a much higher rate than what one would expect to see on an exascale machine in practice. Equally well, the numbers of degrees of freedom n_F that we can simulate is chosen as large as feasible but falls short of what one may expect to achieve on an exascale machine. The laissez-faire mitigation strategy described in Section 2.1 is oblivious of the nature of the fault so long as they are reliably detected. Consequently, our fault simulation procedure does not need to take account of the type of fault. Hence, in order to simulate Algorithm 2, we generate independent realizations for the random matrices involved. In particular,

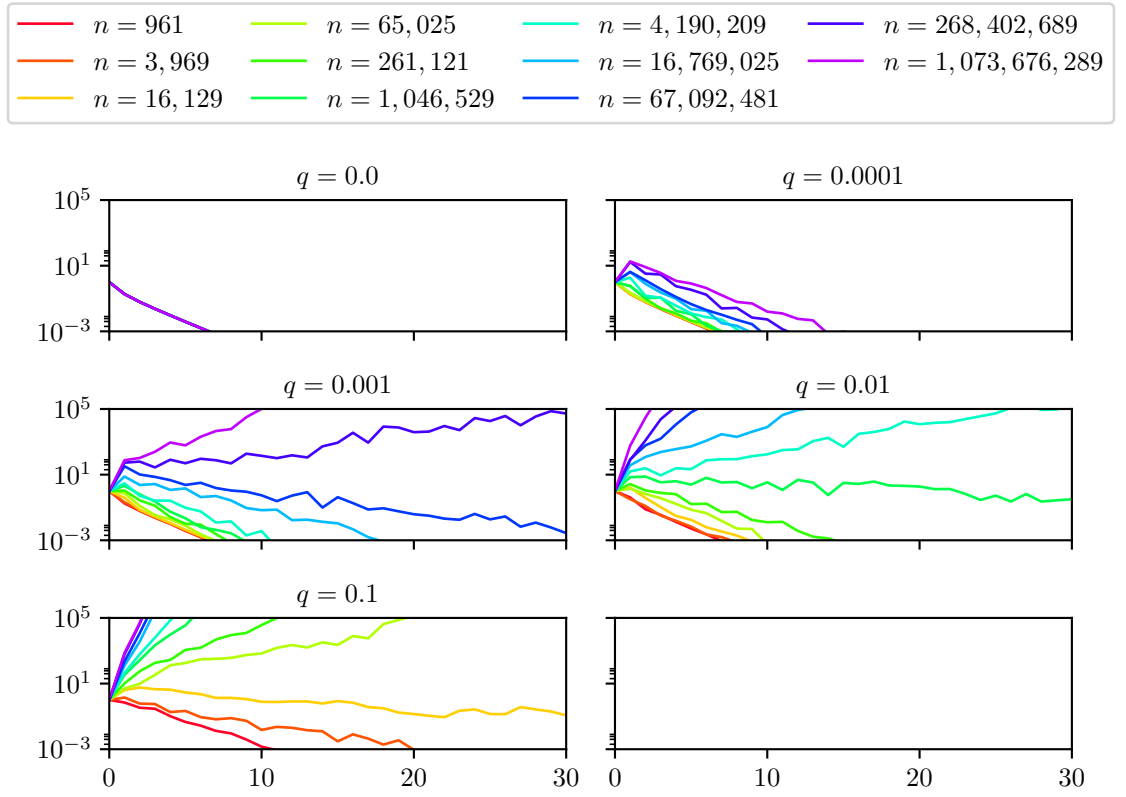


Figure 5.1: Plot of the norm of the residual against iteration number for the Fault-Prone Two Grid Method in the case of discretization of the Poisson problem on a square domain.

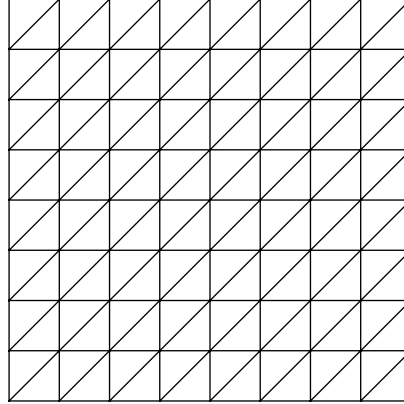


Figure 5.2: Uniform mesh for the square domain.

in order to realize line 5 of Algorithm 2, which contains two fault terms, we sample $n_C + n_F$ independent realizations from an ε -distributed Bernoulli variable (each resulting in a value of zero or one), and multiply residual and coarse-grid right-hand side element by element with the corresponding sampled values. The same approach is applied to simulate faults in lines 4, 9 and 11 of Algorithm 2, where the expression for the optimally damped Jacobi smotherer is given by eq. (2.6).

In Figure 5.1 we plot the norm of the residual after each iteration for a range of values of failure probability ε and system sizes n_F ranging from 1 thousand to 1 billion. These computations were performed on the Titan supercomputer located at Oak Ridge National Laboratory.

It is observed that when $\varepsilon = 0$, the method reaches residual norm 10^{-3} in about 7 iterations uniformly for all values of n_F up to around one billion. However, as ε grows, it is observed that the convergence deteriorates in all cases, with the deterioration being more and more severe as the number of unknowns is increased, until, eventually, the method fails to converge at all.

The fact that the Two Grid Method fails to converge even for such a simple classical problem means one cannot expect a more favorable behavior for realistic practical problems.

In order to verify the scaling law suggested by eq. (5.1) we examine the case

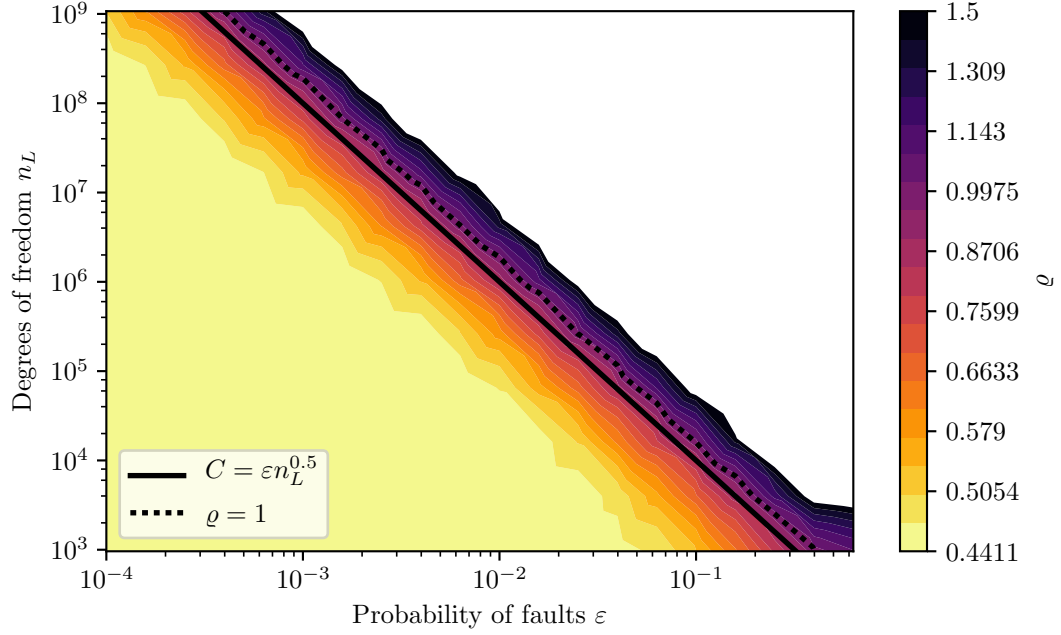


Figure 5.3: Lyapunov spectral radius $\rho(\mathcal{E}^{TG})$ for the iteration matrix for the Fault-Prone Two Grid Method in the case of discretization of the Poisson problem on a square domain.

$d = 2$ further. In particular, we expect convergence to fail when $\varepsilon\sqrt{n_F}$ exceeds some threshold. The Lyapunov spectral radius of the Fault-Prone Two Grid Method is approximated by repeatedly applying Algorithm 2. We perform 1000 iterations for each parameter combination of ε and n_F , and avoid under- and overflow of floating point numbers by renormalizing the iterates. Figure 5.3 shows a contour plot of the variation of the Lyapunov spectral radius with respect to the failure probability ε and the size n_F of the system. The plot indicates the boundary of the region in which the Lyapunov spectral radius exceeds unity, or equally, the region in which the fault-prone iteration no longer converges at all. It will be observed that the lines on which $\varepsilon\sqrt{n_F}$ remains constant (also indicated in the plot) coincide with the contours of the plot in the region in which failure to converge occurs. This supports the scaling law suggested in the estimate eq. (5.1).

5.2 Protecting the Prolongation Restores Resilience

Theorem 12 indicates that the Fault-Prone Two Grid Method will generally not be resilient to faults. What remedial action, in addition to laissez-faire, is needed to restore the uniform convergence of the Fault-Prone Multigrid Method to that of the fault-free scheme? Theorem 18 states that if the *prolongation* operation is protected, meaning that $\mathcal{X}^{(P)} = I$, then the rate of convergence of the Two Grid Method is independent of the number of unknowns and close to the rate of the fault-free method.

Theorem 18.

Let $\mathcal{E}^{TG}(\nu_1, \nu_2) = (\mathcal{E}^{S,post})^{\nu_2} \mathcal{E}^{CG} (\mathcal{E}^{S,pre})^{\nu_1}$ be the iteration matrix of the Fault-Prone Two Grid Method with faults in smoother, residual and restriction, and protected prolongation. Provided Assumptions 1 to 5 hold, and that N and $\mathcal{X}^{(S)}$ commute, we find that

$$\varrho(\mathcal{E}^{TG}(\nu_1, \nu_2)) \leq \min_{\substack{\mu_1 + \mu_2 = \nu_1 + \nu_2 \\ \mu_1, \mu_2 \geq 0}} \|\mathcal{E}^{TG}(\mu_1, \mu_2)\|_2 + C\varepsilon,$$

where the constant C is independent of ε and n_F .

Proof. Adding and subtracting $\mathbb{E}[\mathcal{E}^{TG}(\nu_1, \nu_2)]^{\otimes 2}$ from $\mathbb{E}[(\mathcal{E}^{TG}(\nu_1, \nu_2))^{\otimes 2}]$, we estimate in $\|\bullet\|_{A^2}$

$$\begin{aligned} \rho\left(\mathbb{E}[(\mathcal{E}^{TG}(\nu_1, \nu_2))^{\otimes 2}]\right) &\leq \left\|\mathbb{E}[(\mathcal{E}^{TG}(\nu_1, \nu_2))^{\otimes 2}]\right\|_{A^2} \\ &\leq \left\|\mathbb{E}[\mathcal{E}^{TG}(\nu_1, \nu_2)]\right\|_{A^2}^2 + \left(\left\|\mathbb{E}[\mathcal{E}^{CG}]\right\|_{A^2}^2 + \left\|\mathcal{C}^{(R)}\right\|_{A^2} + \left\|\mathcal{C}^{(\rho)}\right\|_{A^2} + \left\|\mathcal{C}^{(R,\rho)}\right\|_{A^2}\right) \\ &\quad \times \left(\left\|\mathbb{E}[\mathcal{E}^S]\right\|_{A^2}^2 + \left\|\mathcal{C}^{(S)}\right\|_{A^2}\right)^{\nu_1 + \nu_2} \\ &\quad - \left\|\mathbb{E}[\mathcal{E}^{CG}]\right\|_{A^2}^2 \left\|\mathbb{E}[\mathcal{E}^S]\right\|_{A^2}^{2(\nu_1 + \nu_2)}. \end{aligned}$$

We then get by Assumptions 2, 4 and 5 and $(e^\bullet)^2 \leq 1 + 2C\varepsilon + C^2\varepsilon^2 \leq 1 + C\varepsilon$ that

$$\begin{aligned} \left\|\mathcal{C}^{(R)}\right\|_{A^2} &\leq \varepsilon (e^{(R)})^2 \left\|\mathbf{A}_C^{-1} \mathbf{R} \mathbf{A}_F\right\|_2^2 \left\|\mathbf{R}\right\|_2^2 \leq \varepsilon (e^{(R)})^2 \underline{C}_p^2 \overline{C}_p^2 C_A^2 \leq C\varepsilon, \\ \left\|\mathcal{C}^{(\rho)}\right\|_{A^2} &\leq \varepsilon (e^{(\rho)})^2 \left\|\mathbf{P} \mathbf{A}_C^{-1} \mathbf{R} \mathbf{A}_F\right\|_2^2 \leq \varepsilon (e^{(\rho)})^2 C_A^2 \leq C\varepsilon, \end{aligned}$$

$$\left\| \mathbf{C}^{(R,\rho)} \right\|_{A^2} \leq \varepsilon^2 \left\| \mathbf{A}_C^{-1} \mathbf{R} \mathbf{A}_F \right\|_2^2 \left\| \mathbf{R} \right\|_2^2 \leq \varepsilon^2 \underline{C}_p^2 \overline{C}_p^2 C_A^2 \leq C \varepsilon^2.$$

We also estimate

$$\left\| \mathbf{C}^{(S)} \right\|_{A^2} = \left\| \mathbf{A}_F^{\otimes 2} \mathbf{V}^{(S)} \mathbf{N}^{\otimes 2} \right\|_2 = \left\| (\mathbf{A}_F \mathbf{N})^{\otimes 2} \mathbf{V}^{(S)} \right\|_2 \leq \varepsilon \left\| \mathbf{A}_F \mathbf{N} \right\|_2^2 \leq C \varepsilon.$$

Here, we used Assumptions 3 and 5 and that $\mathbf{N}$ and $\mathbf{X}^{(S)}$ commute. Moreover,

$$\begin{aligned} \left\| \mathbb{E} [\mathbf{E}^{CG}] \right\|_{A^2} &= \left\| \mathbb{E} [\mathbf{E}^{CG}]^T \right\|_2 = \left\| \mathbb{E} [\mathbf{E}^{CG}] \right\|_2 \\ &\leq \left\| \mathbf{E}^{CG} \right\|_2 + |1 - e^{(R)} e^{(\rho)}| \left\| \mathbf{I} - \mathbf{E}^{CG} \right\|_2 \\ &\leq C_A (1 + |1 - e^{(R)} e^{(\rho)}|) \leq C (1 + \varepsilon), \\ \left\| \mathbb{E} [\mathbf{E}^S] \right\|_{A^2} &= \left\| \mathbb{E} [\mathbf{E}^S]^T \right\|_2 = \left\| \mathbb{E} [\mathbf{E}^S] \right\|_2 \\ &\leq \left\| \mathbf{E}^S \right\|_2 + |1 - e^{(S)}| \left\| \mathbf{I} - \mathbf{E}^S \right\|_2 \leq C(1 + \varepsilon), \\ \left\| \mathbb{E} [\mathbf{E}^{TG}(\nu_1, \nu_2)] \right\|_{A^2} &= \left\| \mathbb{E} [\mathbf{E}^{TG}(\nu_1, \nu_2)]^T \right\|_2 = \left\| \mathbb{E} [\mathbf{E}^{TG}(\nu_2, \nu_1)] \right\|_2 \\ &= \left\| \mathbf{E}^{TG}(\nu_2, \nu_1) \right\|_2 + \left\| \mathbb{E} [\mathbf{E}^{TG}(\nu_2, \nu_1)] \right\|_2 - \left\| \mathbf{E}^{TG}(\nu_2, \nu_1) \right\|_2 \\ &\leq \left\| \mathbf{E}^{TG}(\nu_2, \nu_1) \right\|_2 + C \varepsilon. \end{aligned}$$

Here, we used that by Assumption 5

$$|1 - e^{(\bullet)}| \leq C \varepsilon, \quad (e^{(\bullet)})^2 \leq 1 + C \varepsilon, \quad |1 - e^{(R)} e^{(\rho)}| \leq C \varepsilon.$$

Collecting all the terms, we have

$$\rho \left(\mathbb{E} \left[(\mathbf{E}^{TG}(\nu_1, \nu_2))^{\otimes 2} \right] \right) \leq \left\| \mathbf{E}^{TG}(\nu_2, \nu_1) \right\|_2^2 + C \varepsilon.$$

so that we finally obtain

$$\varrho \left(\mathbf{E}^{TG}(\nu_1, \nu_2) \right) \leq \left\| \mathbf{E}^{TG}(\nu_2, \nu_1) \right\|_2 + C \varepsilon$$

We conclude by observing that the Lyapunov spectral radius, just as the ordinary spectral radius, is invariant with respect to cyclic permutations. $\square$

Theorem 18 makes no additional assumptions about the origin of the problem or the solver hierarchy, and allows for more general fault patterns and smootheners than Theorem 12 and as such, is more generally applicable.

We apply the Fault-Prone Two Grid Method incorporating protection of the prolongation to the same test problem as above. In Figure 5.5 we plot the evolution of the residual norm for varying system sizes and fault rates, and in Figure 5.4 we plot the Lyapunov spectral radius of the iteration matrix. Observe that the rate of convergence is independent of the number of unknowns and even stays below unity (indicating convergence) for fault probabilities as high as $\varepsilon \approx 0.6$. Moreover, Figure 5.5 seems to indicate that while the asymptotic behavior is independent of the number of unknowns, the first iteration is not.

In practice, we are faced with the problem of protecting the prolongation. Some developers have proposed architectures on which certain components are less susceptible to faults, and the prolongation operator would be a candidate for this treatment. We shall present a practical approach to the protection of the prolongation in Chapter 7.

5.3 Protecting the Restriction Fails to Give Resilience

The previous result raises the question as to whether protection of any other component of the algorithm can give similar (or better) results. From a theoretical aspect, this does not seem to be the case, as careful examination of the proof of Theorem 18 reveals. A provably resilient method can only be obtained if either the prolongation or both restriction and residual in the coarse-grid correction are protected. To confirm whether or not this proves to be the case in practice, we repeat the previous simulations, using instead a protected restriction operation. The measured rates of convergence as shown in Figure 5.6 resemble the behavior in the case without protection. Another possible alternative might be to protect the calculation

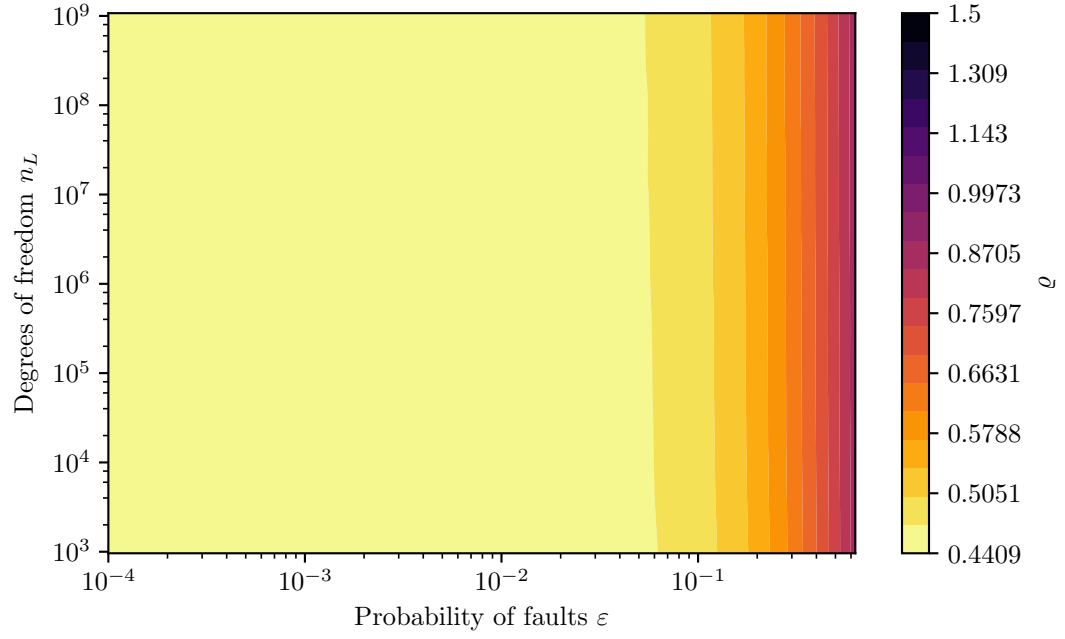


Figure 5.4: Lyapunov spectral radius $\rho(\mathcal{E}^{TG})$ for the iteration matrix for the Fault-Prone Two Grid Method with protected prolongation in the case of discretization of Poisson problem on a square domain.

of the residual. However, this is computationally unattractive, since the residual constitutes typically the most expensive part of the method and will therefore have a larger overhead than protection of the prolongation or restriction.

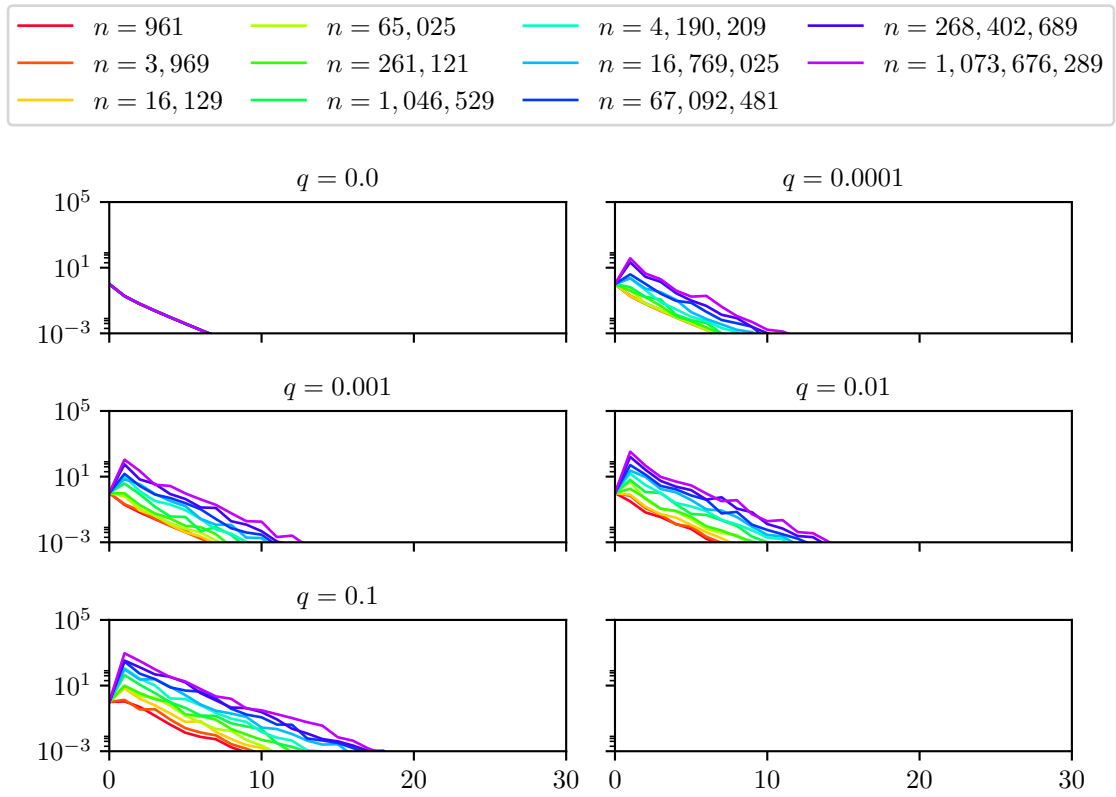


Figure 5.5: Plot of the norm of the residual against iteration number for the Fault-Prone Two Grid Method with protected prolongation in the case of discretization of Poisson problem on square domain.

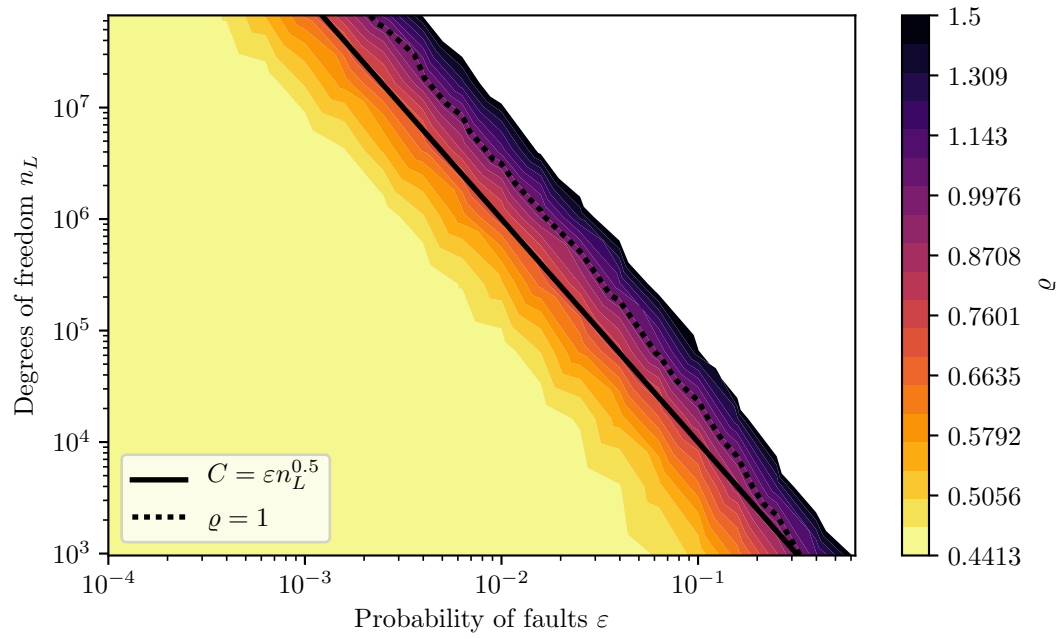


Figure 5.6: Lyapunov spectral radius $\rho(\mathcal{E}^{TG})$ for the iteration matrix for the Fault-Prone Two Grid Method with protected restriction in the case of discretization of Poisson problem on a square domain.

Chapter 6

Analysis of the Fault-Prone Multigrid Method

While the two grid method is, as a solver, mostly of academic interest, the behavior of the multigrid method under the impact of faults is of great practical importance. We extend the result of Theorem 18 to the W-cycle. This theorem mirrors the classical implication of W-cycle convergence by two grid convergence outlined in Section 1.3, but applies to the Lyapunov spectral radius needed for the analysis of the Fault-Prone Multigrid Method. No additional assumptions are required beyond these needed for the classical multigrid analysis.

Theorem 19.

Provided Assumptions 1 to 5 hold, that the prolongation is protected, that the number of smoothing steps is sufficient, and that ε is sufficiently small, the fault-prone multigrid method converges with a rate bounded by

$$\varrho(\mathcal{E}_\ell(\nu_1, \nu_2, \gamma)) \leq \begin{cases} \frac{\gamma}{\gamma-1} C_{TG} + C\varepsilon, & \gamma \geq 2, \\ \frac{2}{1+\sqrt{1-4C_*C_{TG}}} C_{TG} + C\varepsilon, & \gamma = 2, \end{cases}$$

where C is independent of ε and ℓ and

$$C_{TG} = \max_{\ell} \|\mathbf{E}_\ell^{TG}(\nu_2, \nu_1)\|_2 \leq C_S C_A \eta(\nu_2),$$

$$C_* = C_S \underline{C}_p \overline{C}_p (C_S + C_A \eta(\nu_2)).$$

Setting the fault rate $\varepsilon = 0$ in the above bounds recovers the classical estimates from Section 1.3 for the fault-free multigrid method, since $\varrho(\mathcal{E}_\ell)$ reduces to $\rho(\mathbf{E}_\ell)$ for $\varepsilon = 0$. Just as the Two Grid result, this Theorem makes no assumptions about the origin of the solver hierarchy, and does not rely on a particular choice of smoothener. The proof of Theorem 19 requires several steps that are recorded in the results below.

Throughout this section, C will be a generic constant whose value can change from line to line, but which is independent of ℓ and ε .

We set

$$\mathbb{E} [\mathcal{X}_\ell^{(\bullet)}] = e_\ell^{(\bullet)} \mathbf{I}, \quad \mathbb{V} [\mathcal{X}_\ell^{(\bullet)}] = \mathbf{V}_\ell^{(\bullet)}.$$

and expand

$$\begin{aligned} \mathbb{E} [(\mathcal{E}_\ell^S)^{\otimes 2}] &= \mathbb{E} [\mathcal{E}_\ell^S]^{\otimes 2} + \mathbf{C}_\ell^{(S)}, \\ \mathbf{C}_\ell^{(S)} &= \mathbf{V}_\ell^{(S)} (\mathbf{N}_\ell \mathbf{A}_\ell)^{\otimes 2}, \end{aligned}$$

using Lemma 9. Similar to the fault-free case, we derive recursive inequalities for the iteration matrix of the multigrid method.

Theorem 20.

If Assumptions 1 to 5 hold, then the iteration matrix of the Fault-Prone Multigrid Method satisfies the following recursive inequalities:

$$\|\mathbb{E} [\mathcal{E}_\ell(\nu_1, \nu_2, \gamma)]\|_{A^2} \leq \|\mathbb{E} [\mathcal{E}_\ell^{TG}(\nu_1, \nu_2)]\|_{A^2} + C\varepsilon + C_* \|\mathbb{E} [\mathcal{E}_{l-1}(\nu_1, \nu_2, \gamma)]\|_{A^2}^\gamma, \quad (6.1)$$

and

$$\begin{aligned} \|\mathbb{E} [\mathcal{E}_\ell(\nu_1, \nu_2, \gamma)^{\otimes 2}]\|_{A^2} &\leq \|\mathbb{E} [\mathcal{E}_\ell^{TG}(\nu_1, \nu_2)^{\otimes 2}]\|_{A^2} + C\varepsilon \\ &\quad + 2C_* \|\mathbb{E} [\mathcal{E}_\ell^{TG}(\nu_1, \nu_2)]\|_{A^2} \|\mathbb{E} [\mathcal{E}_{l-1}(\nu_1, \nu_2, \gamma)]\|_{A^2}^\gamma \end{aligned} \quad (6.2)$$

$$+ C_*^2 \left\| \mathbb{E} [\boldsymbol{\mathcal{E}}_{l-1} (\nu_1, \nu_2, \gamma)^{\otimes 2}] \right\|_{A^2}^\gamma,$$

with

$$C_* = C_S \underline{C}_p \overline{C}_p (C_S + C_A \eta(\nu_2)).$$

The constants C and C_* are independent of the level ℓ and ε .

If $\varepsilon = 0$, the second inequality matches exactly the recursive inequality that we would use for the fault-free case, while the first inequality matches the tensor square power of the same inequality.

Proof. Since the proof on level ℓ only involves the levels ℓ and $\ell - 1$, we will drop the first subscript and replace the second one with a subscript C .

We can write the iteration matrix as

$$\begin{aligned} \boldsymbol{\mathcal{E}} (\nu_1, \nu_2, \gamma) &= (\boldsymbol{\mathcal{E}}^{S, \text{post}})^{\nu_2} \left(I - P (I - (\boldsymbol{\mathcal{E}}_C)^\gamma) \mathbf{A}_C^{-1} \boldsymbol{\mathcal{X}}_C^{(R)} \mathbf{R} \boldsymbol{\mathcal{X}}^{(\rho)} \mathbf{A} \right) (\boldsymbol{\mathcal{E}}^{S, \text{pre}})^{\nu_1} \\ &= \underbrace{(\boldsymbol{\mathcal{E}}^{S, \text{post}})^{\nu_2} \left(I - P \mathbf{A}_C^{-1} \boldsymbol{\mathcal{X}}_C^{(R)} \mathbf{R} \boldsymbol{\mathcal{X}}^{(\rho)} \mathbf{A} \right) (\boldsymbol{\mathcal{E}}^{S, \text{pre}})^{\nu_1}}_{\boldsymbol{\mathcal{E}}^{TG}(\nu_1, \nu_2)} \\ &\quad + \underbrace{(\boldsymbol{\mathcal{E}}^{S, \text{post}})^{\nu_2} P (\boldsymbol{\mathcal{E}}_C)^\gamma \mathbf{A}_C^{-1} \boldsymbol{\mathcal{X}}_C^{(R)} \mathbf{R} \boldsymbol{\mathcal{X}}^{(\rho)} \mathbf{A} (\boldsymbol{\mathcal{E}}^{S, \text{pre}})^{\nu_1}}_{=: \boldsymbol{\mathcal{C}}}. \end{aligned}$$

Since

$$\begin{aligned} \left\| \mathbb{E} [(\boldsymbol{\mathcal{E}}^S)^{\nu_2} P] \right\|_{A^2} &= \left\| \mathbf{A} \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_2} P \mathbf{A}_C^{-1} \right\|_2 \\ &= \left\| \mathbf{A}_C^{-1} \mathbf{R} \mathbf{A} \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_2} \right\|_2 \\ &\leq \underline{C}_p (\left\| \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_2} \right\|_2 + \left\| \mathbf{E}^{CG} \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_2} \right\|_2) \\ &\leq \underline{C}_p (C_S + C_A \eta(\nu_2)) + C\varepsilon, \\ \left\| \mathbb{E} [\mathbf{A}_C^{-1} \boldsymbol{\mathcal{X}}_C^{(R)} \mathbf{R} \boldsymbol{\mathcal{X}}^{(\rho)} \mathbf{A}] \right\|_{A^2} &\leq \left| e_C^{(R)} e^{(\rho)} \right| \left\| \mathbf{R} \right\|_2 \leq \overline{C}_p + C\varepsilon, \\ \left\| \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_1} \right\|_{A^2} &\leq C_S + C\varepsilon, \end{aligned}$$

we have by Assumptions 4 and 5 that

$$\left\| \mathbb{E} [\boldsymbol{\mathcal{C}}] \right\|_{A^2} \leq \left\| \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_2} P \right\|_{A^2} \left\| \mathbb{E} [\boldsymbol{\mathcal{E}}_C] \right\|_{A^2}^\gamma \left\| \mathbb{E} [\mathbf{A}_C^{-1} \boldsymbol{\mathcal{X}}_C^{(R)} \mathbf{R} \boldsymbol{\mathcal{X}}^{(\rho)} \mathbf{A}] \right\|_{A^2} \left\| \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_1} \right\|_{A^2}$$

$$\leq C_* \|\mathbb{E}[\boldsymbol{\mathcal{E}}_C]\|_{A^2}^\gamma + C\varepsilon. \quad (6.3)$$

Therefore the first inequality of the theorem follows from

$$\|\mathbb{E}[\boldsymbol{\mathcal{E}}]\|_{A^2} \leq \|\mathbb{E}[\boldsymbol{\mathcal{E}}^{TG}]\|_{A^2} + \|\mathbb{E}[\boldsymbol{\mathcal{C}}]\|_{A^2}.$$

The second inequality is more work, but follows the same path:

$$\|\mathbb{E}[\boldsymbol{\mathcal{E}}^{\otimes 2}]\|_{A^2} \leq \left\| \mathbb{E}[(\boldsymbol{\mathcal{E}}^{TG})^{\otimes 2}] \right\|_{A^2} + \|\mathbb{E}[\boldsymbol{\mathcal{C}}^{\otimes 2}]\|_{A^2} + 2 \|\mathbb{E}[\boldsymbol{\mathcal{E}}^{TG} \otimes \boldsymbol{\mathcal{C}}]\|_{A^2}. \quad (6.4)$$

We start by bounding the second term on the right-hand side.

$$\begin{aligned} \|\mathbb{E}[\boldsymbol{\mathcal{C}}^{\otimes 2}]\|_{A^2} &\leq \left\| \mathbb{E}[(\boldsymbol{\mathcal{E}}^S)^{\nu_2} \boldsymbol{P}]^{\otimes 2} \right\|_{A^2} \|\mathbb{E}[\boldsymbol{\mathcal{E}}_C^{\otimes 2}]\|_{A^2}^\gamma \\ &\quad \times \left\| \mathbb{E}[(\boldsymbol{A}_C^{-1} \boldsymbol{\mathcal{X}}_C^{(R)} \boldsymbol{R} \boldsymbol{\mathcal{X}}^{(\rho)} \boldsymbol{A})^{\otimes 2}] \right\|_{A^2} \left\| \mathbb{E}[(\boldsymbol{\mathcal{E}}^S)^{\nu_1}]^{\otimes 2} \right\|_{A^2} \end{aligned} \quad (6.5)$$

We have by

$$\begin{aligned} \left\| \mathbb{E}[(\boldsymbol{\mathcal{E}}^S)^{\nu_2} \boldsymbol{P}]^{\otimes 2} \right\|_{A^2} &\leq \left\| (\mathbb{E}[\boldsymbol{\mathcal{E}}^S]^{\nu_2} \boldsymbol{P})^{\otimes 2} \right\|_{A^2} + C \|\boldsymbol{C}^{(S)}\|_{A^2} \\ &= \|\boldsymbol{A} \mathbb{E}[\boldsymbol{\mathcal{E}}^S]^{\nu_2} \boldsymbol{P} \boldsymbol{A}_C^{-1}\|_2^2 + C\varepsilon \\ &= \|\boldsymbol{A}_C^{-1} \boldsymbol{R} \boldsymbol{A} \mathbb{E}[\boldsymbol{\mathcal{E}}^S]^{\nu_2}\|_2^2 + C\varepsilon \\ &\leq \underline{C}_p^2 (\|\mathbb{E}[\boldsymbol{\mathcal{E}}^S]^{\nu_2}\|_2 + \|\boldsymbol{E}^{CG} \mathbb{E}[\boldsymbol{\mathcal{E}}^S]^{\nu_2}\|_2)^2 + C\varepsilon \\ &\leq \underline{C}_p^2 (C_S + C_A \eta(\nu_2))^2 + C\varepsilon. \end{aligned} \quad (6.6)$$

Since by Assumptions 4 and 5

$$\begin{aligned} \left\| \mathbb{E}[(\boldsymbol{A}_C^{-1} \boldsymbol{\mathcal{X}}_C^{(R)} \boldsymbol{R} \boldsymbol{\mathcal{X}}^{(\rho)} \boldsymbol{A})^{\otimes 2}] \right\|_{A^2} &= \left\| \mathbb{E}[(\boldsymbol{\mathcal{X}}_C^{(R)} \boldsymbol{R} \boldsymbol{\mathcal{X}}^{(\rho)})^{\otimes 2}] \right\|_2 \\ &\leq \left\| \mathbb{E}[(\boldsymbol{\mathcal{X}}_C^{(R)})^{\otimes 2}] \right\|_2 \|\boldsymbol{R}\|_2^2 \left\| \mathbb{E}[(\boldsymbol{\mathcal{X}}^{(\rho)})^{\otimes 2}] \right\|_2 \\ &\leq \left(\|\mathbb{V}[\boldsymbol{\mathcal{X}}_C^{(R)}]\|_2 + \|\mathbb{E}[\boldsymbol{\mathcal{X}}_C^{(R)}]\|_2^2 \right) \|\boldsymbol{R}\|_2^2 \\ &\quad \times \left(\|\mathbb{V}[\boldsymbol{\mathcal{X}}^{(\rho)}]\|_2 + \|\mathbb{E}[\boldsymbol{\mathcal{X}}^{(\rho)}]\|_2^2 \right) \\ &\leq (\varepsilon + (e^{(R)})^2) (\varepsilon + (e^{(\rho)})^2) \|\boldsymbol{R}\|_2^2 \end{aligned} \quad (6.7)$$

$$\leq \overline{C}_p^2 + C\varepsilon.$$

and

$$\begin{aligned} \left\| \mathbb{E} \left[((\boldsymbol{\mathcal{E}}^S)^{\nu_1})^{\otimes 2} \right] \right\|_{A^2} &= \left\| \left(\mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\otimes 2} + \boldsymbol{C}^{(S)} \right)^{\nu_1} \right\|_{A^2} \\ &\leq \left\| \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_1} \right\|_{A^2}^2 + C\varepsilon \leq C_S^2 + C\varepsilon, \end{aligned} \quad (6.8)$$

Inserting eqs. (6.8), (6.6) and (6.7) into eq. (6.5), we have

$$\left\| \mathbb{E} [\boldsymbol{C}^{\otimes 2}] \right\|_{A^2} \leq C_*^2 \left\| \mathbb{E} [\boldsymbol{\mathcal{E}}_C^{\otimes 2}] \right\|_{A^2}^\gamma + C\varepsilon. \quad (6.9)$$

Now we turn to the third term of eq. (6.4). We split

$$\begin{aligned} \boldsymbol{\mathcal{E}}^{TG} &= (\boldsymbol{\mathcal{E}}^{S,\text{post}})^{\nu_2} \left(\boldsymbol{I} - \boldsymbol{P} \boldsymbol{A}_C^{-1} \boldsymbol{\mathcal{X}}_C^{(R)} \boldsymbol{R} \boldsymbol{\mathcal{X}}^{(\rho)} \boldsymbol{A} \right) (\boldsymbol{\mathcal{E}}^{S,\text{pre}})^{\nu_1} \\ &= \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_2} \boldsymbol{E}^{CG} \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_1} \\ &\quad + ((\boldsymbol{\mathcal{E}}^{S,\text{post}})^{\nu_2} - \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_2}) \boldsymbol{E}^{CG} \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_1} \\ &\quad + \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_2} \boldsymbol{E}^{CG} ((\boldsymbol{\mathcal{E}}^{S,\text{pre}})^{\nu_1} - \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_1}) \\ &\quad + ((\boldsymbol{\mathcal{E}}^{S,\text{post}})^{\nu_2} - \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_2}) \boldsymbol{E}^{CG} ((\boldsymbol{\mathcal{E}}^{S,\text{pre}})^{\nu_1} - \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_1}) \\ &\quad + (\boldsymbol{\mathcal{E}}^{S,\text{post}})^{\nu_2} \boldsymbol{P} \left(\boldsymbol{A}_C^{-1} \boldsymbol{R} \boldsymbol{A} - \boldsymbol{A}_C^{-1} \boldsymbol{\mathcal{X}}_C^{(R)} \boldsymbol{R} \boldsymbol{\mathcal{X}}^{(\rho)} \boldsymbol{A} \right) (\boldsymbol{\mathcal{E}}^{S,\text{pre}})^{\nu_1} \end{aligned}$$

so that

$$\begin{aligned} &\left\| \mathbb{E} [\boldsymbol{\mathcal{E}}^{TG} \otimes \boldsymbol{C}] \right\|_{A^2} \\ &\leq \left\| \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_2} \boldsymbol{E}^{CG} \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_1} \right\|_{A^2} \left\| \mathbb{E} [\boldsymbol{C}] \right\|_{A^2} \\ &\quad + \left\| \boldsymbol{E}^{CG} \right\|_{A^2} \left\| \boldsymbol{P} \right\|_{A^2} \left\| \mathbb{E} [\boldsymbol{\mathcal{E}}_C] \right\|_{A^2}^\gamma \left\| \mathbb{E} \left[\boldsymbol{A}_C^{-1} \boldsymbol{\mathcal{X}}_C^{(R)} \boldsymbol{R} \boldsymbol{\mathcal{X}}^{(\rho)} \boldsymbol{A} \right] \right\|_{A^2} \\ &\quad \left\{ \left\| \mathbb{E} [(\boldsymbol{\mathcal{E}}^S)^{\nu_2} \otimes ((\boldsymbol{\mathcal{E}}^S)^{\nu_2} - \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_2})] \right\|_{A^2} \left\| \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_1} \right\|_{A^2}^2 \right. \\ &\quad + \left\| \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_2} \right\|_{A^2}^2 \left\| \mathbb{E} [(\boldsymbol{\mathcal{E}}^S)^{\nu_1} \otimes ((\boldsymbol{\mathcal{E}}^S)^{\nu_1} - \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_1})] \right\|_{A^2} \\ &\quad + \left\| \mathbb{E} [(\boldsymbol{\mathcal{E}}^S)^{\nu_2} \otimes ((\boldsymbol{\mathcal{E}}^S)^{\nu_2} - \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_2})] \right\|_{A^2} \\ &\quad \left. \times \left\| \mathbb{E} [(\boldsymbol{\mathcal{E}}^S)^{\nu_1} \otimes ((\boldsymbol{\mathcal{E}}^S)^{\nu_1} - \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_1})] \right\|_{A^2} \right\} \\ &\quad + \left\| \mathbb{E} \left[(\boldsymbol{A}_C^{-1} \boldsymbol{R} \boldsymbol{A}) \otimes (\boldsymbol{A}_C^{-1} \boldsymbol{\mathcal{X}}_C^{(R)} \boldsymbol{R} \boldsymbol{\mathcal{X}}^{(\rho)} \boldsymbol{A}) - (\boldsymbol{A}_C^{-1} \boldsymbol{\mathcal{X}}_C^{(R)} \boldsymbol{R} \boldsymbol{\mathcal{X}}^{(\rho)} \boldsymbol{A})^{\otimes 2} \right] \right\|_{A^2} \end{aligned} \quad (6.10)$$

$$\times \|\mathbb{E}[\boldsymbol{\mathcal{E}}_C]\|_{A^2}^\gamma \left\| \mathbb{E} \left[((\boldsymbol{\mathcal{E}}^S)^\nu \boldsymbol{P})^{\otimes 2} \right] \right\|_{A^2} \left\| \mathbb{E} \left[((\boldsymbol{\mathcal{E}}^S)^{\nu_1})^{\otimes 2} \right] \right\|_{A^2}$$

Now, we have by Assumptions 2 and 4 that

$$\|\boldsymbol{E}^{CG}\|_{A^2} = \|\boldsymbol{E}^{CG}\|_2 \leq C_A, \quad (6.11)$$

$$\|\boldsymbol{P}\|_{A^2} = \|\boldsymbol{A} \boldsymbol{P} \boldsymbol{A}_C^{-1}\|_2 \leq \underline{C}_p, \quad (6.12)$$

and by Assumptions 4 and 5 that

$$\begin{aligned} & \left\| \mathbb{E} \left[(\boldsymbol{A}_C^{-1} \boldsymbol{R} \boldsymbol{A}) \otimes (\boldsymbol{A}_C^{-1} \boldsymbol{\mathcal{X}}_C^{(R)} \boldsymbol{R} \boldsymbol{\mathcal{X}}^{(\rho)} \boldsymbol{A}) - (\boldsymbol{A}_C^{-1} \boldsymbol{\mathcal{X}}_C^{(R)} \boldsymbol{R} \boldsymbol{\mathcal{X}}^{(\rho)} \boldsymbol{A})^{\otimes 2} \right] \right\|_{A^2} \\ &= \left\| \mathbb{E} \left[\boldsymbol{R} \otimes (\boldsymbol{\mathcal{X}}_C^{(R)} \boldsymbol{R} \boldsymbol{\mathcal{X}}^{(\rho)}) - (\boldsymbol{\mathcal{X}}_C^{(R)} \boldsymbol{R} \boldsymbol{\mathcal{X}}^{(\rho)})^{\otimes 2} \right] \right\|_2 \\ &\leq \|\boldsymbol{R}\|_2^2 \left\{ \left| e_C^{(R)} e^{(\rho)} \right| \left| 1 - e_C^{(R)} e^{(\rho)} \right| + (e^{(\rho)})^2 \left\| \mathbb{V} [\boldsymbol{\mathcal{X}}_C^{(R)}] \right\|_2 \right. \\ &\quad \left. + (e_C^{(R)})^2 \left\| \mathbb{V} [\boldsymbol{\mathcal{X}}^{(\rho)}] \right\|_2 + \left\| \mathbb{V} [\boldsymbol{\mathcal{X}}_C^{(R)}] \right\|_2 \left\| \mathbb{V} [\boldsymbol{\mathcal{X}}^{(\rho)}] \right\|_2 \right\} \\ &\leq C\varepsilon \end{aligned} \quad (6.13)$$

and that

$$\begin{aligned} & \left\| \mathbb{E} [(\boldsymbol{\mathcal{E}}^S)^\nu \otimes ((\boldsymbol{\mathcal{E}}^S)^\nu - \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^\nu)] \right\|_{A^2} \\ &= \left\| \mathbb{E} [(\boldsymbol{\mathcal{E}}^S)^{\otimes 2}]^\nu - \left(\mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\otimes 2} \right)^\nu \right\|_{A^2} \\ &= \left\| \left(\mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\otimes 2} + \boldsymbol{C}^{(S)} \right)^\nu - \left(\mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\otimes 2} \right)^\nu \right\|_{A^2} \\ &\leq C \left\| \boldsymbol{C}^{(S)} \right\|_{A^2} \leq C\varepsilon. \end{aligned} \quad (6.14)$$

Also,

$$\begin{aligned} & \left\| \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_2} \boldsymbol{E}^{CG} \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_1} \right\|_{A^2} \\ &\leq \left\| \mathbb{E} [\boldsymbol{\mathcal{E}}^{TG}(\nu_1, \nu_2)] \right\|_{A^2} + \left| 1 - e_C^{(R)} \right| \left| 1 - e^{(\rho)} \right| \left\| \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_1} (\boldsymbol{I} - \boldsymbol{E}^{CG}) \mathbb{E} [\boldsymbol{\mathcal{E}}^S]^{\nu_2} \right\|_2 \\ &= \left\| \mathbb{E} [\boldsymbol{\mathcal{E}}^{TG}(\nu_1, \nu_2)] \right\|_{A^2} + C\varepsilon. \end{aligned} \quad (6.15)$$

Hence we find by inserting eqs. (6.3), (6.11), (6.12) and (6.13) to (6.14) into eq. (6.10)

$$\left\| \mathbb{E} [\boldsymbol{\mathcal{E}}^{TG} \otimes \boldsymbol{C}] \right\|_{A^2} \leq C_* \left\| \mathbb{E} [\boldsymbol{\mathcal{E}}^{TG}] \right\|_{A^2} \|\mathbb{E} [\boldsymbol{\mathcal{E}}_C]\|_{A^2}^\gamma + C\varepsilon. \quad (6.16)$$

Therefore, the second recursive inequality

$$\begin{aligned} \|\mathbb{E} [\boldsymbol{\varepsilon}^{\otimes 2}]\|_{A^2} &\leq \left\| \mathbb{E} \left[(\boldsymbol{\varepsilon}^{TG})^{\otimes 2} \right] \right\|_{A^2} + C\varepsilon \\ &\quad + 2C_* \|\mathbb{E} [\boldsymbol{\varepsilon}^{TG}]\|_{A^2} \|\mathbb{E} [\boldsymbol{\varepsilon}_C]\|_{A^2}^\gamma \\ &\quad + (C_*)^2 \|\mathbb{E} [\boldsymbol{\varepsilon}_C^{\otimes 2}]\|_{A^2}^\gamma \end{aligned}$$

follows from eqs. (6.4), (6.9) and (6.16). $\square$

We adapt Lemma 2 to the tensor space setting:

Lemma 21.

Let the elements of the sequences $\{\eta_\ell\}_{\ell \geq 1}$ and $\{\eta_{\ell, \otimes 2}\}_{\ell \geq 1}$ be bounded as follows

$$\begin{aligned} \eta_1 &\leq \xi, & \eta_\ell &\leq \xi + C_* \eta_{\ell-1}^\gamma, \quad \ell \geq 2, \\ \eta_{1, \otimes 2}^2 &\leq \xi_{\otimes 2}^2, & \eta_{\ell, \otimes 2}^2 &\leq \xi_{\otimes 2}^2 + 2\xi C_* \eta_{\ell-1}^\gamma + C_*^2 \eta_{\ell-1, \otimes 2}^{2\gamma}, \quad \ell \geq 2. \end{aligned}$$

If $\gamma \geq 2$, $C_* \gamma > 1$ and $\max \{\xi, \xi_{\otimes 2}\} \leq \frac{\gamma-1}{\gamma} \frac{1}{(C_* \gamma)^{\frac{1}{\gamma-1}}}$, then

$$\max \{\eta_\ell, \eta_{\ell, \otimes 2}\} \leq \frac{\gamma}{\gamma-1} \max \{\xi, \xi_{\otimes 2}\} < 1, \quad \ell \geq 1.$$

In the case of $\gamma = 2$, we have

$$\max \{\eta_\ell, \eta_{\ell, \otimes 2}\} \leq \frac{2 \max \{\xi, \xi_{\otimes 2}\}}{1 + \sqrt{1 - 4C_* \max \{\xi, \xi_{\otimes 2}\}}} < 1, \quad \ell \geq 1.$$

Proof. Set $\alpha_\ell := \max \{\eta_\ell, \eta_{\ell, \otimes 2}\}$ and $\beta := \max \{\xi, \xi_{\otimes 2}\}$. Then

$$\begin{aligned} \eta_1 &\leq \beta, & \eta_\ell &\leq \beta + C_* \alpha_{\ell-1}^\gamma, \quad \ell \geq 2, \\ \eta_{1, \otimes 2}^2 &\leq \beta^2, & \eta_{\ell, \otimes 2}^2 &\leq \beta^2 + 2\beta C_* \alpha_{\ell-1}^\gamma + C_*^2 \alpha_{\ell-1}^{2\gamma}, \quad \ell \geq 2, \end{aligned}$$

and hence

$$\alpha_1 \leq \beta, \quad \alpha_\ell \leq \beta + C_* \alpha_{\ell-1}^\gamma, \quad \ell \geq 2.$$

Application of lemma 2 gives for $\gamma \geq 2$

$$\alpha_\ell \leq \frac{\gamma}{\gamma-1} \beta < 1$$

and for $\gamma = 2$

$$\alpha_\ell \leq \frac{2\beta}{1 + \sqrt{1 - 4C_*\beta}} < 1.$$

□

This finally permits us to give prove the main result of this chapter:

Proof of Theorem 19. Set

$$\begin{aligned}\eta_\ell &:= \|\mathbb{E}[\boldsymbol{\mathcal{E}}_\ell(\nu_1, \nu_2, \gamma)]\|_{A^2}, \\ \eta_{\ell, \otimes 2} &:= \|\mathbb{E}[\boldsymbol{\mathcal{E}}_\ell(\nu_1, \nu_2, \gamma)^{\otimes 2}]\|_{A^2}^{\frac{1}{2}}, \\ \xi &:= \max_\ell \|\boldsymbol{E}_\ell^{TG}(\nu_2, \nu_1)\|_2 + C\varepsilon, \\ \xi_{\otimes 2} &:= \max_\ell \|\boldsymbol{E}_\ell^{TG}(\nu_2, \nu_1)\|_2 + C\varepsilon.\end{aligned}$$

It follows from

$$\begin{aligned}\mathbb{E}[\boldsymbol{\mathcal{E}}_\ell^S] &= \boldsymbol{E}_\ell^S + \left(1 - e_\ell^{(S)}\right) (\boldsymbol{I} - \boldsymbol{E}_\ell^S), \\ \mathbb{E}[\boldsymbol{\mathcal{E}}_\ell^{CG}] &= \boldsymbol{E}_\ell^{CG} + \left(1 - e_{\ell-1}^{(R)} e_\ell^{(\rho)}\right) (\boldsymbol{I} - \boldsymbol{E}_\ell^{CG}),\end{aligned}$$

and Assumptions 2, 3 and 5 that

$$\|\mathbb{E}[\boldsymbol{\mathcal{E}}_\ell^{TG}(\nu_1, \nu_2)]\|_{A^2} \leq \|\boldsymbol{E}_\ell^{TG}(\nu_2, \nu_1)\|_2 + C\varepsilon,$$

and therefore with eq. (6.1) that $\eta_\ell \leq \xi + C_*\eta_{\ell-1}^\gamma$. By Theorem 18 and eq. (6.2), we find $\eta_{\ell, \otimes 2}^2 \leq \xi_{\otimes 2}^2 + 2\xi C_*\eta_{\ell-1}^\gamma + C_*^2\eta_{\ell-1, \otimes 2}^2$. Since C_* , ξ and $\xi_{\otimes 2}$ decrease as the number of smoothening steps increases, Lemma 21 can be applied for a sufficient number of smoothening steps. □

6.1 The Effect of Protection of the Prolongation

Theorem 19 assumes that the prolongation operator is protected. We investigate whether this assumption can be relaxed by considering a numerical example where

we do not protect the prolongation. Specifically, we consider the Poisson equation in two dimensions

$$\begin{cases} -\Delta u = f & \text{in } \Omega = [0, 1]^2, \\ u = 0 & \text{on } \partial\Omega, \end{cases} \quad (6.17)$$

and in order to rule out extraneous effects, we use a uniform mesh, a discretization with piecewise linear finite elements, and optimally damped Jacobi pre- and post-smootheners. In particular, it is well established that the fault-free multigrid method converges in this setting. Assumptions 1 to 4 are satisfied, as for example shown in [62]. On every level, residual, restriction, prolongation and smootheners are subject to componentwise faults, as given in eq. (2.2). We consider problems of size between 1 million and 1 billion degrees of freedom, and fault probabilities between 10^{-4} and 0.6. To minimize floating point contamination in the approximation of the Lyapunov spectral radius, the right-hand side f is taken to be zero, a non-zero random initial iterate is chosen, and after each iteration the current iterate is renormalized. In Figure 6.1, we plot the evolution of the residual norm with respect to the iteration number for the case of laissez-faire mitigation in prolongation, restriction, residual and smootheners. We can see that the number of degrees of freedom n_L adversely affects the rate of convergence, even leading to divergence for large number of unknowns. Estimates of the Lyapunov spectral radius are obtained as the geometric average of 1000 iterations of Fault-Prone Multigrid, and are displayed in Section 6.1.

In the Two Grid setting which was discussed in Chapter 5, we found that the asymptotic rate for this problem scales like $\sqrt{n_L \varepsilon^2}$. Similarly, as seen in Section 6.1, we find

$$\varrho(\mathcal{E}_L) = C_0 + \mathcal{O}\left(\sqrt{n_L \varepsilon^2}\right) \quad (6.18)$$

for the multilevel case, with C_0 being a constant related to the fault-free method. This means that the method is not fault resilient, since for any given rate of faults

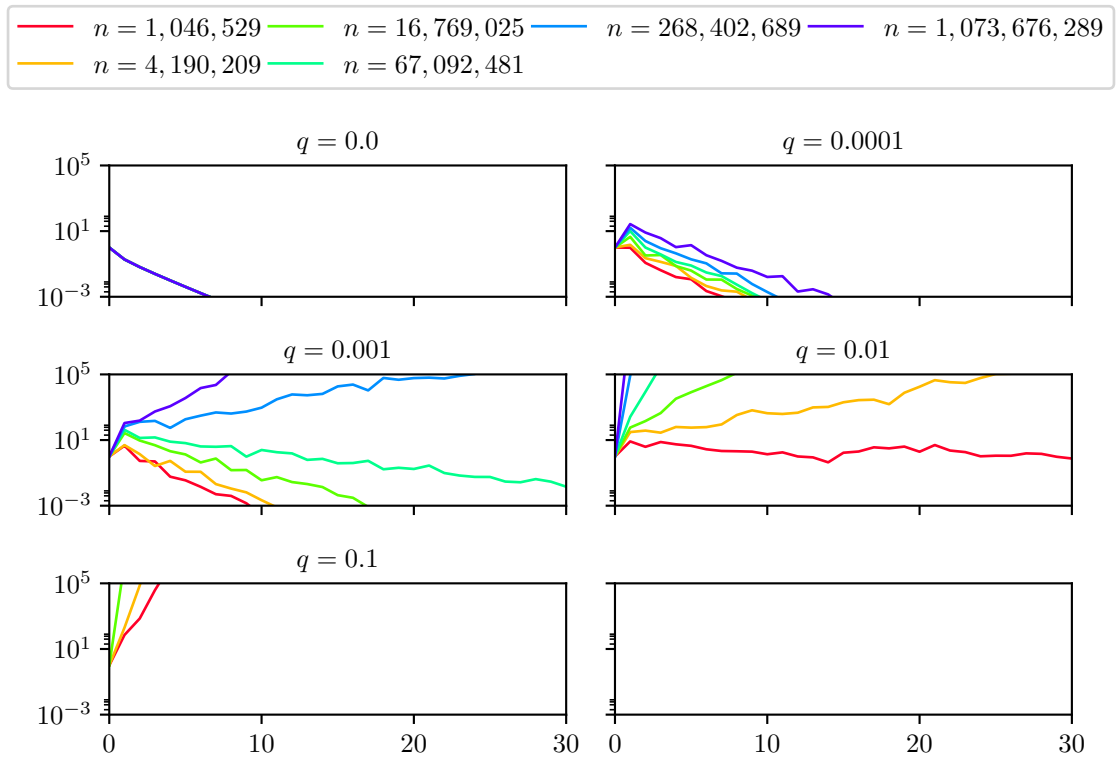


Figure 6.1: Plots of the norm of the residual against iteration number for the Fault-Prone W-Cycle Multigrid Method in the case of discretization of Poisson problem on square domain.

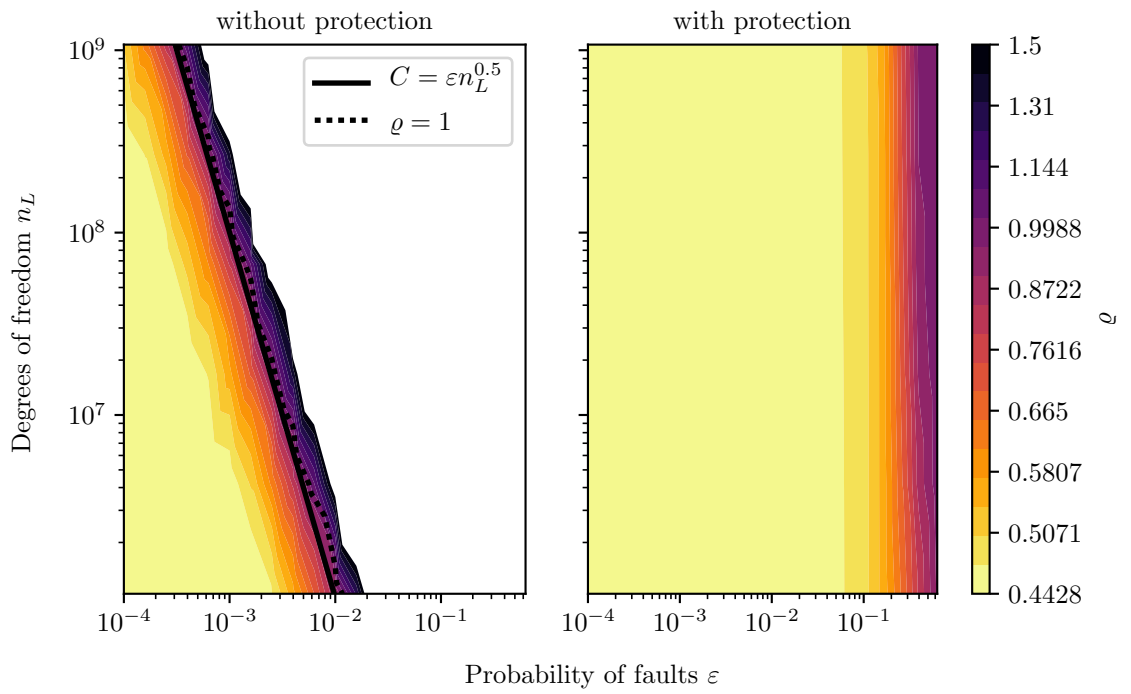


Figure 6.2: Lyapunov spectral radius $\rho(\mathcal{E}_L)$ for the iteration matrix for the Fault-Prone W-Cycle Multigrid method in the case of discretization of Poisson problem on a square domain. Without protected prolongation (*left*), and with protected prolongation (*right*).

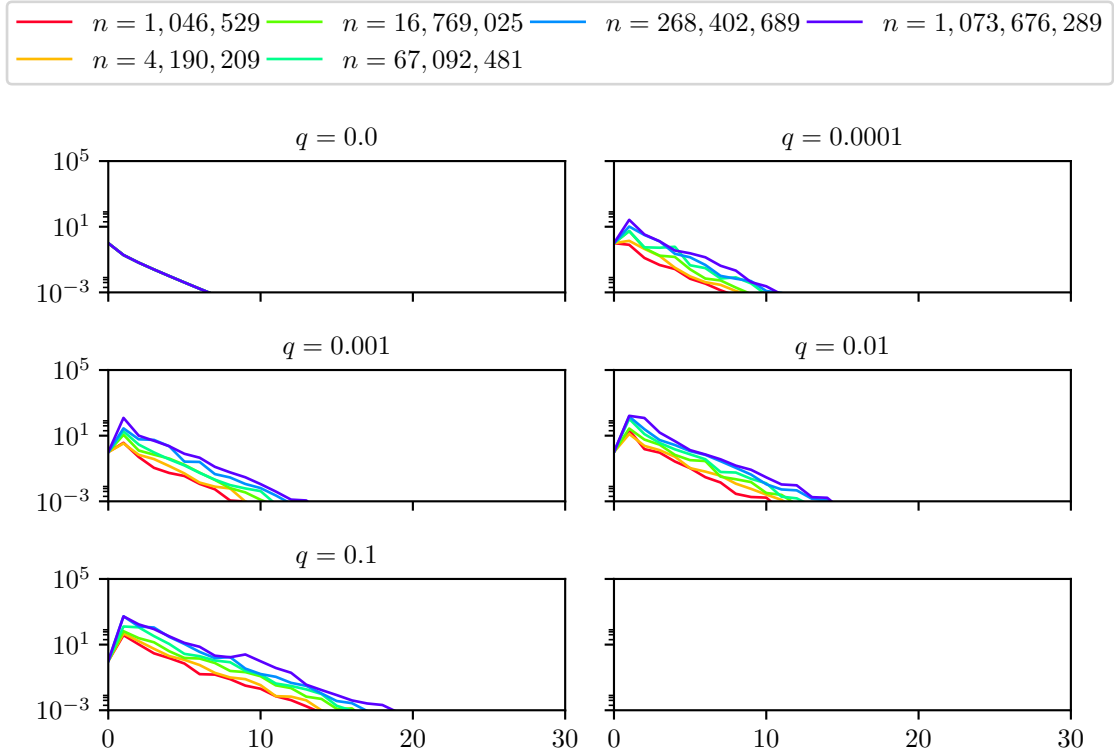


Figure 6.3: Plot of the norm of the residual against iteration number for the Fault-Prone W-Cycle Multigrid Method with protected prolongation in the case of discretization of Poisson problem on square domain.

ε , there exists a maximum problem size above which multigrid diverges. Therefore additional protection is mandatory in the multilevel case as well.

We repeat the same experiments, this time with a protected prolongation, i.e. $\mathcal{X}_\ell^{(P)} = \mathbf{I}$. This produces the desired independence of the problem size, as predicted by Theorem 19 and as shown by the evolution of the residual in Figure 6.3 and by the estimated rate of convergence in Section 6.1. We further illustrate the results with several test problems that pose more of a challenge to the multigrid solver.

6.2 Adaptively Refined Meshes

A second numerical example illustrates the results of Theorem 19 for the case of an adaptively refined mesh. We consider the 2D magnetostatics problem for a three phase 6/4 switched reluctance motor as depicted in Figure 6.4. Gauss's law and Ampère's law are given by

$$\operatorname{div} \vec{B} = 0, \quad \operatorname{curl} \vec{H} = J,$$

where $\vec{B}$ is the magnetic flux density, $\vec{H}$ the magnetic field intensity and J the current density. $\vec{B}$ and $\vec{H}$ are linked through the constitutive relation $\vec{B} = \mu \vec{H}$ with magnetic permeability μ . Using the magnetic vector potential $\vec{B} = \operatorname{curl} u$, the system can be rewritten as

$$-\frac{1}{\mu} \Delta u = J.$$

This gives rise to a variational problem with

$$a(u, v) = \int_{\Omega} \frac{1}{\mu} \operatorname{grad} u \cdot \operatorname{grad} v, \quad F(v) = \int_{\Omega} J v.$$

The permeability μ is 5200 in the rotor and the stator, and 1 everywhere else. The current density J is unity in the coils and 0 everywhere else. Using continuous piecewise linear finite elements and classical residual-based local a posteriori error indicators [50], we adaptively refine an initial mesh shown in Figure 6.4. We use Red-Green refinement [12] coupled with a Dörfler marking strategy [43] with refinement parameter 0.5. The solver hierarchy is generated without injection of faults; then the problem is solved for componentwise faults and varying fault rates.

Section 6.2 depicts the approximate rate of convergence of the W-Cycle with one step of Jacobi pre- and post-smoothing each and without protection of the prolongation. We can see that although neither the W-Cycle nor the adaptively refined mesh are covered by the two grid result from Chapter 5, the convergence estimate is qualitatively correct. Added protection of the prolongation operation

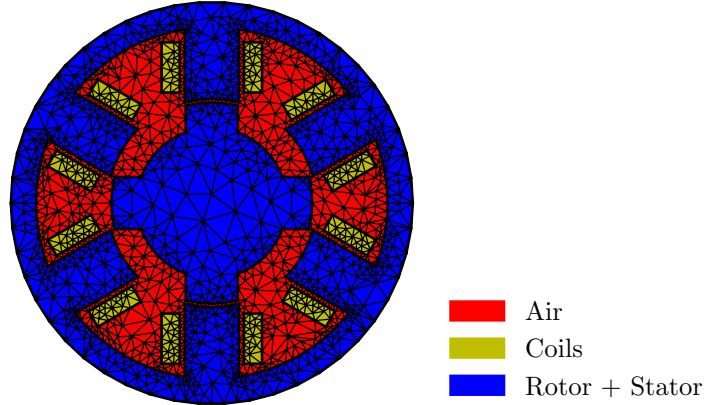


Figure 6.4: Initial mesh for the motor problem.

leads to a fault resilient method, as shown in Section 6.2. Small variations with respect to the number of degrees of freedom are due to varying coarsening ratios of the levels. The results match Theorem 19, even though Assumption 2 is not satisfied.

6.3 Higher Spatial Dimension and Higher Order Finite Elements

We now demonstrate that fault resilience of the W-cycle is retained for a 3D partial differential equation and higher order (quadratic) continuous finite elements, and solve the Poisson equation on the Fichera cube.

$$\begin{aligned} -\Delta u &= f & \text{in } \Omega &= [0, 2]^3 \setminus [1, 2]^3, \\ u &= u_0 & \text{on } \partial\Omega \end{aligned}$$

The geometry and its uniform meshing are shown in Figure 6.6. The system is partitioned using METIS [71] into blocks of size 2^{14} and distributed over the compute nodes. We inject nodewise faults as given by eq. (2.3). The problem sizes considered range from 1.7 million to about 940 million degrees of freedom. In Section 6.3, we show the approximate rate of convergence as well as a level set of the

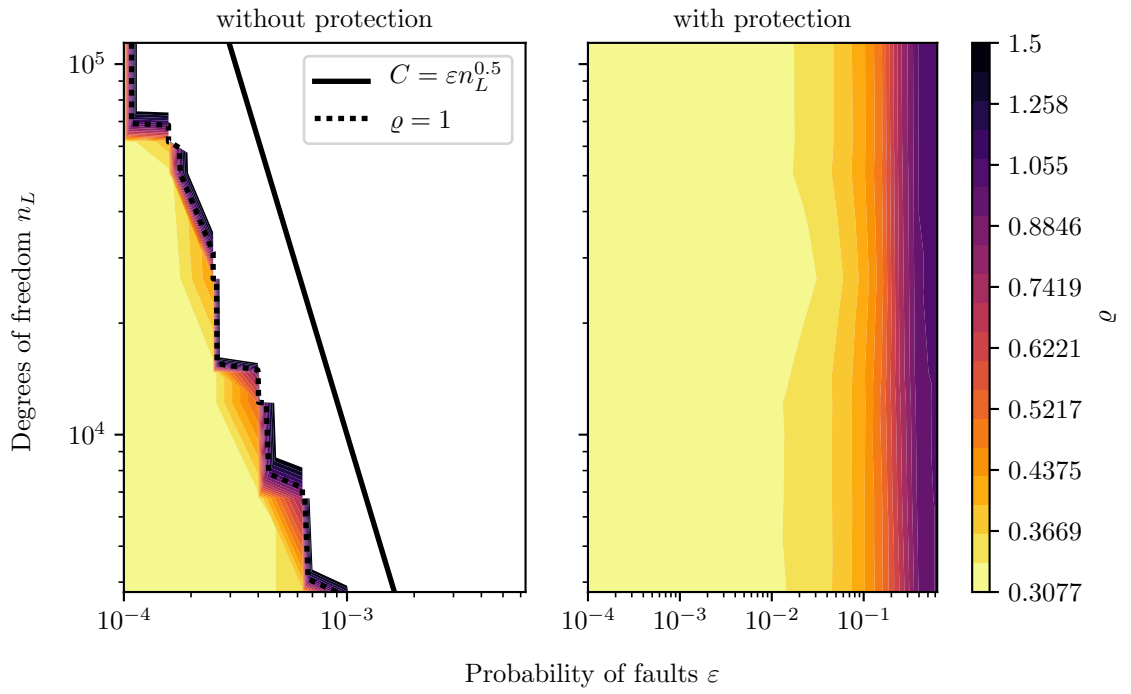


Figure 6.5: Lyapunov spectral radius $\varrho(\mathcal{E}_L)$ for the iteration matrix for the Fault-Prone W-Cycle Multigrid Method in the case of discretization of the motor problem. Without protected prolongation (*left*), and with protected prolongation (*right*).

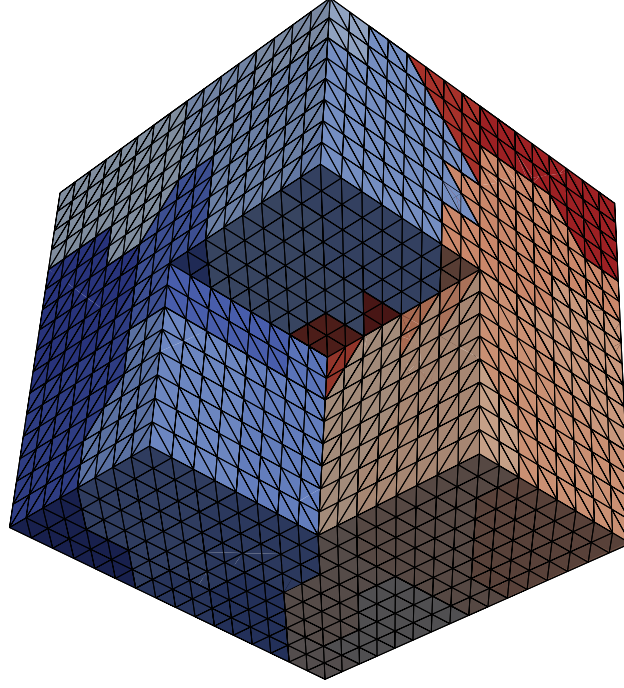


Figure 6.6: Mesh for the Fichera cube, indicating partitioning from METIS.

error bound obtained for the Fault-Prone Two Grid Method in Chapter 5. It can be seen that the rate of convergence in the multilevel case behaves like

$$\varrho(\mathcal{E}_L) = C_0 + \varepsilon \sqrt[6]{n_L},$$

where C_0 is a constant that is related to the fault-free method. We notice that the onset of divergent behavior indeed happens at a slower rate as compared to the 2D setting. Protection of the prolongation again makes the method fault resilient, as shown in Section 6.3. The choice of higher order elements plays no significant role in the convergence behavior of the method.

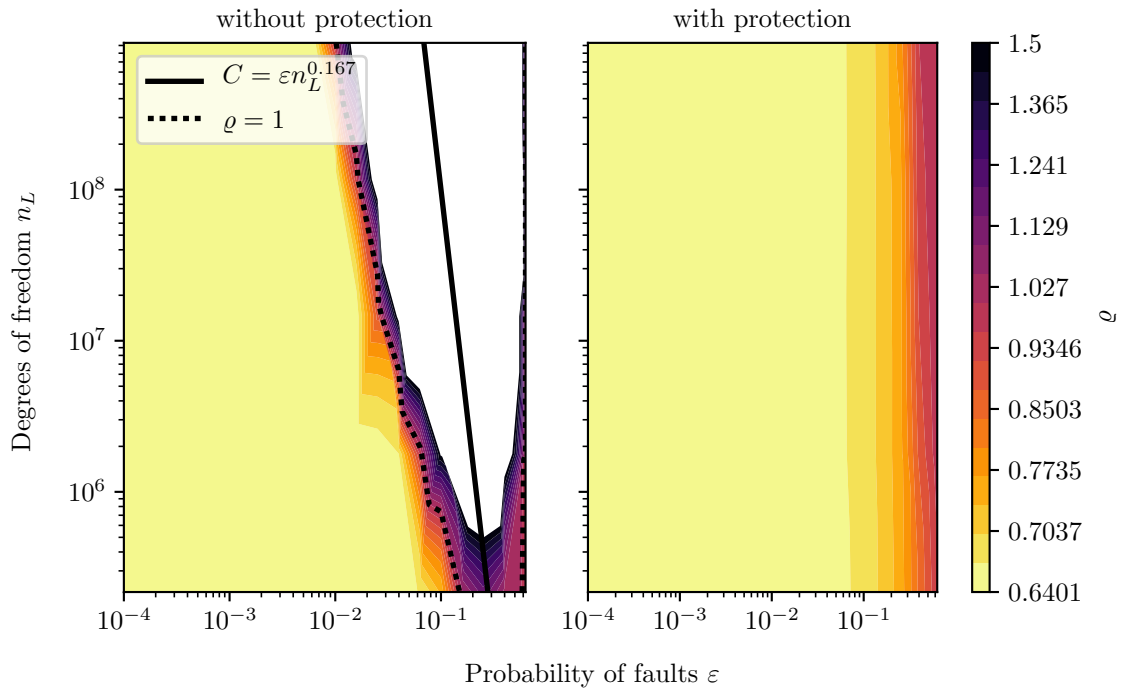


Figure 6.7: Lyapunov spectral radius $\varrho(\mathcal{E}_L)$ for the iteration matrix for the Fault-Prone W-Cycle Multigrid Method in the case of discretization of the Fichera cube. Without protected prolongation (*left*), and with protected prolongation (*right*).

Chapter 7

Implementation Issues

In the previous chapters, we assumed that faults can be perfectly detected. In practice, faults cannot be perfectly detected, nor is it possible to perfectly protect the prolongation. The question therefore arises of whether Theorem 19 has any practical relevance. In Section 7.1, we present one possible simple approach for fault detection and show that the behavior of Fault-Prone Multigrid found in eq. (6.18) is recovered. In Section 7.2 we turn to the issue of protecting the prolongation operator.

7.1 Detection of Soft Faults

The laissez-faire fault mitigation strategy requires fault detection. While this is straight-forward in the case of hard faults, soft faults are more problematic. Various techniques have been suggested [89]. Here, we present a simple approach based on replication [64] in which a fault-prone component is repeated K times ($K \geq 2$) and the results are compared for consistency. Algorithm 5 shows how the strategy is used in conjunction with laissez-faire in the detection of faults in the computation of a generic matrix-vector product $y = Mx$. Instead of the action of M , only its fault-prone equivalent $\mathcal{M}$ is available in practice, where $\mathcal{M} \in \mathbb{R}^{n \times n}$ is a random

matrix.

Algorithm 5 Detection of faults using K replicas and laissez-faire mitigation for the fault-prone matrix-vector product $y = \mathcal{M}x$.

```

1: for  $i \leftarrow 1$  to  $n$  do
2:   for  $j \leftarrow 1$  to  $K$  do
3:      $w_j \leftarrow (\mathcal{M}x)_i$            //  $j$ -th replica
4:   if  $w_1 = w_2 = \dots = w_K$  and  $|w_1| < 10^{16}$  then
5:      $y_i \leftarrow w_1$            // value accepted
6:   else
7:      $y_i \leftarrow 0$            // Laissez-faire

```

The basic idea behind Algorithm 5 is to declare an operation as being fault-free if all replicas are in agreement, otherwise a fault is deemed to have occurred and laissez-faire is triggered. This means that there are three possible outcomes for the i -th component of the output y from Algorithm 5:

- (C) the correct value of $y_i = (Mx)_i$ is returned,
- (M) an error is detected, the laissez-faire mitigation is triggered, and $y_i = 0$,
- (U) a fault occurs but remains undetected, and y_i is a corrupted value.

Our analysis in Chapter 6 caters for the cases (C) and (M), but does not take account of case (U). Suppose that the probability of a replica being corrupted is $\varepsilon \ll 1$. Then the probability of *all* K replicas being corrupted in *exactly* the same way is $\mathcal{O}(\varepsilon^K)$, showing that the likelihood of (U) occurring is exponentially small in K . Hence, the probabilities of the outcomes of Algorithm 5 are

$$\mathbb{P}(C) = 1 - \mathcal{O}(\varepsilon),$$

$$\mathbb{P}(M) = \mathcal{O}(\varepsilon),$$

$$\mathbb{P}(U) = \mathcal{O}(\varepsilon^K).$$

We illustrate the effect of using our ad-hoc fault detection strategy given in Algorithm 5 in Figure 7.1 for the Poisson problem eq. (6.17). In particular, the faults

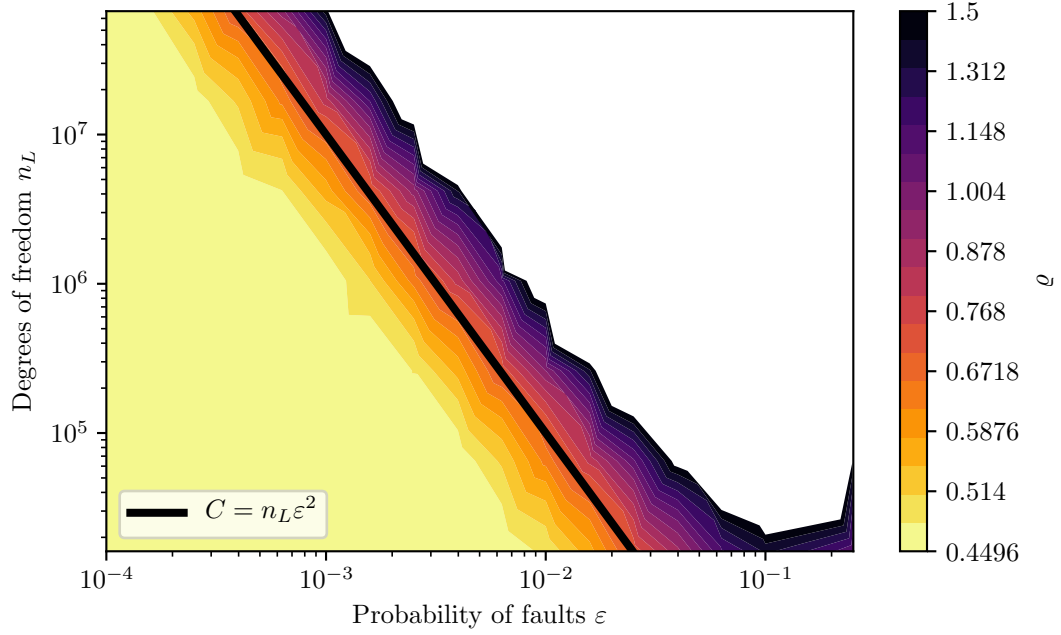


Figure 7.1: Lyapunov spectral radius $\varrho(\mathcal{E}_L)$ for the iteration matrix for the Fault-Prone W-Cycle Multigrid Method in the case of discretization of the 2D Poisson problem with fault detection using $K = 2$ replicas in every operation.

are modeled using bit flipping in which any bit in the floating point representation is flipped with a small but non-zero probability such that the overall probability of a floating point number being corrupted is ε . The results are given in Figure 7.1 in the simplest case where $K = 2$ replicas are used in Algorithm 5 and all operations. It is observed that, even in this simplest case, the convergence behavior mirrors that which would be obtained with perfect detection.

7.2 Protection of the Prolongation

Theorem 19 requires that the prolongation operations

$$x_\ell \leftarrow x_\ell + P e_{\ell-1}$$

are computed exactly, i.e. without *any* faults. This is clearly not possible in practice. We have seen that in the absence of full protection of the prolongation, the Fault-Prone Multigrid converges at a rate given by eq. (6.18), where ε is the underlying failure rate of the machine. Moreover, Theorem 19 suggests that the $\mathcal{O}(\sqrt{n_L \varepsilon^2})$ term is entirely due to faults in the prolongation (with the faults coming purely from other components of multigrid contributing with higher order terms). At any rate, it is clear that if we can enhance the reliability of the prolongation *sufficiently* (without being necessarily exact), then one can expect to ameliorate the factor $\mathcal{O}(\sqrt{n_L \varepsilon^2})$ in the bound of the Lyapunov spectral radius, e.g. by obtaining a growth of $\mathcal{O}(n_L \varepsilon^\alpha)$ for some $\alpha > 2$.

Is it possible to improve the likelihood of detecting and mitigating faults in the prolongation beyond $\mathbb{P}(M) = \mathcal{O}(\varepsilon)$? Algorithm 5 may be regarded as being overly conservative. For instance, if all but one of the replicas are in agreement, then intuitively it seems likely that the majority are correct. Algorithm 6 implements the approach suggested by this argument by looking for agreement amongst a subset of k_P of the K_P replicas for the prolongation. This added freedom alters the likelihood

Algorithm 6 Protection of the fault-prone prolongation $x_\ell \leftarrow x_\ell + \mathcal{P}e_{\ell-1}$ with up to K_P replicas, acceptance threshold k_P and laissez-faire mitigation.

```

1: for  $i \leftarrow 1$  to  $n_\ell$  do
2:   for  $j \leftarrow 1$  to  $K_P$  do
3:      $w_j \leftarrow (\mathcal{P}e_{\ell-1})_i$            //  $j$ -th replica
4:     if  $k_P$  replicas have matching value  $w$  and  $|w| < 10^{16}$  then
5:        $(x_\ell)_i \leftarrow (x_\ell)_i + w$        // value accepted
6:       break

```

at which the three events occur in the prolongation:

$$\mathbb{P}(C_P) = 1 - \mathcal{O}(\varepsilon^{K_P - k_P + 1}),$$

$$\mathbb{P}(M_P) = \mathcal{O}(\varepsilon^{K_P - k_P + 1}),$$

$$\mathbb{P}(U_P) = \mathcal{O}(\varepsilon^{k_P}).$$

In particular, when Algorithm 6 is applied to the computation of the prolongation

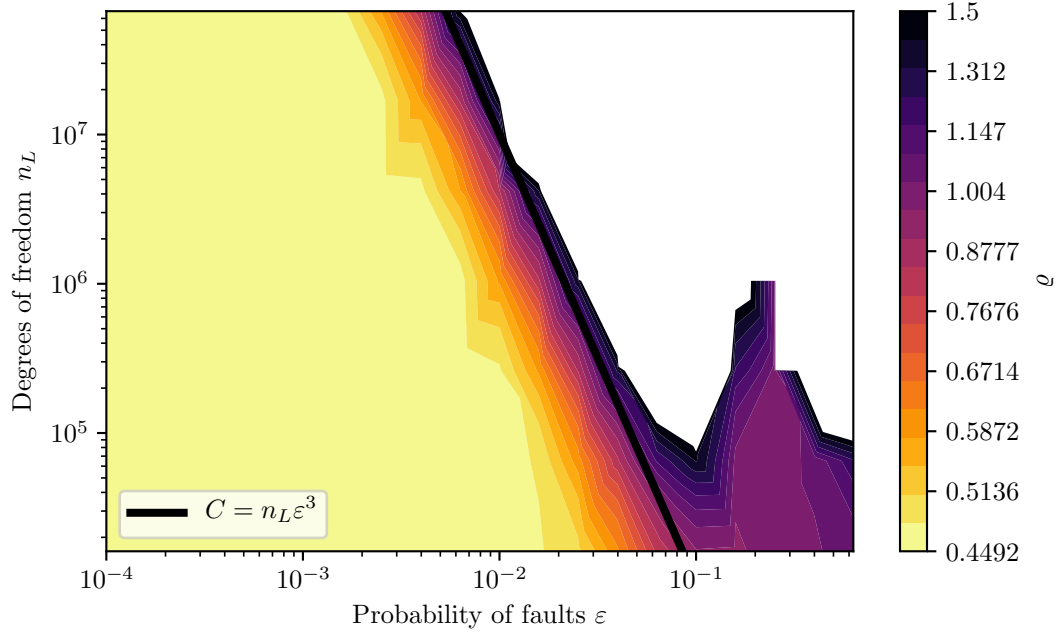


Figure 7.2: Lyapunov spectral radius $\varrho(\mathcal{E}_L)$ for the iteration matrix for the Fault-Prone W-Cycle Multigrid Method with $K_P = 4$ replicas and acceptance threshold $k_P = 3$ in the prolongation and $K = 3$ replicas in all other operations.

with parameters $k_P = 3$ and $K_P = 4$, we obtain an overall rate of $\mathbb{P}(M_P) = \mathcal{O}(\varepsilon^2)$ at which mitigation occurs. Figure 7.2 shows the results obtained by applying Algorithm 6 with $k_P = 3$ and $K_P = 4$ to the prolongation, and Algorithm 5 with $K = 3$ to all other operations. It is seen that the strategy has resulted in the factor $\mathcal{O}(\sqrt{n_L \varepsilon^2})$ in the Lyapunov spectral radius being replaced by $\mathcal{O}(n_L \varepsilon^\alpha)$ with $\alpha = 3$.

Would a cheaper detection and protection strategy (e.g. with $K_P = 3$, $k_P = 2$ and $K = 2$) suffice? Figure 7.3 shows the results obtained employing a variety of different choices for K , K_P and k_P . It is seen that the cheapest strategy that results in the growth $\mathcal{O}(\sqrt{n_L \varepsilon^2})$ being replaced by $\mathcal{O}(n_L \varepsilon^\alpha)$, $\alpha > 2$, is indeed $K_P = 4$, $k_P = 3$ and $K = 3$.

Moreover, if one wishes to obtain a given rate $\alpha \geq 2$, then one can choose $K = k_P = \alpha$ and $K_P = 2\alpha - 2$. As problems in three spatial dimensions are more resilient to the effect of laissez-faire mitigation, as reflected by Theorem 12 and the

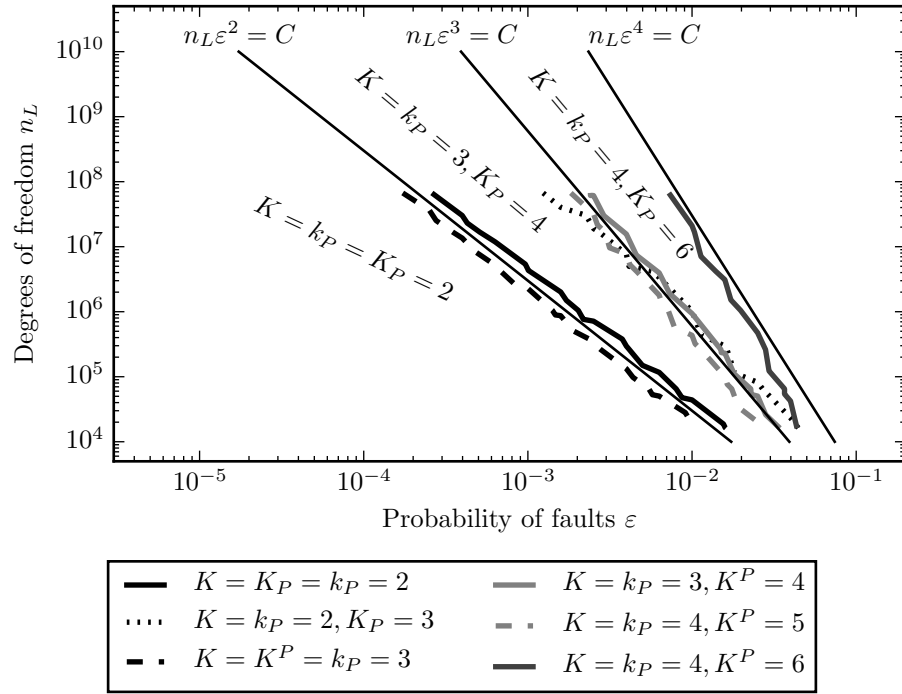


Figure 7.3: Level sets $\varrho(\mathcal{E}_L) = 0.57$ (50% more iterations than the ideal fault-free case) for different levels of detection and protection in the case of the 2D Poisson problem. Optimal parameter combinations are marked in thick solid lines. The obtained optimal detection and protection parameters for each region in the ε - n_L -plane are marked in the plot.

numerical evidence in Chapter 6, we expect the strategy $K = K_P = k_P = \alpha$ to be sufficient for $2 \leq \alpha \leq 6$.

Chapter 8

Fault-Prone Domain Decomposition Methods

A second class of widely used iterative linear solvers are so-called domain decomposition methods. The main idea behind domain decomposition methods is to divide up the problem domain into several potentially overlapping subdomains, and solve the system restricted to every subdomain concurrently. Therefore, a domain decomposition method is similar in spirit to a block Jacobi smoothener. In typical applications, the subdomain size is chosen so that each compute node in the parallel computation can fit one subdomain problem. In order to obtain a scalable method, a coarse grid solve needs to be incorporated for global exchange of information. While multigrid methods achieve optimal complexity by coarsening the problem only mildly and using inexpensive solvers on each level, domain decomposition methods coarsen aggressively, so that typically only two levels are used. In turn, the considered problem on the different levels are significantly more expensive to solve [41, 88, 94].

Domain decomposition methods can be divided up into two categories: solvers without and with overlap of the subdomains. In this chapter, we introduce the later kind, and discuss an extension of the framework for the analysis of *laissez-faire* fault

mitigation to blockwise faults in one and two level domain decomposition solvers. In particular, we explain the effect on faults in unknowns that are shared by several subdomains.

8.1 Overlapping Domain Decomposition Methods

Let the polygonal domain $\Omega \subset \mathbb{R}^d$ be divided up into a shape regular, quasi-uniform partition $\mathcal{T}_h$ of simplices of size h . Partition $\mathcal{T}_h$ into N parts $\mathcal{T}_h^{(i,0)}$, $i = 1, \dots, N$, corresponding to subdomains $\Omega^{(i,0)}$. The extended triangulations $\mathcal{T}_h^{(i,\delta)}$, $\delta = 1, 2, \dots$ are obtained by adding δ exterior layers of cells to $\mathcal{T}_h^{(i,0)}$. We denote the corresponding overlapping subdomains by $\Omega^{(i,\delta)}$. Let $V := H_0^1(\Omega)$ and V_h be the space of piecewise linears on $\mathcal{T}_h$, and $V_h^{(i,\delta)}$ the space of piecewise linears on $\mathcal{T}_h^{(i,\delta)}$, with respective dimensions $n = \dim V_h$ and $n^{(i,\delta)} = \dim V_h^{(i,\delta)}$. Let ϕ_j the nodal basis functions of V_h , and let $\vec{\phi}$ and $\vec{\phi}^{(i,\delta)}$ be the vectors of basis functions in V_h and $V_h^{(i,\delta)}$ respectively, with index sets $\mathcal{I}$ and $\mathcal{I}^{(i,\delta)} \subset \mathcal{I}$. Moreover, let $z_j \in \Omega$ be the node corresponding to the basis function ϕ_j , $j \in \mathcal{I}$. The canonical injection

$$\mathcal{P}^{(i,\delta)} : V_h^{(i,\delta)} \hookrightarrow V_h$$

gives rise to the prolongation matrix $\mathbf{P}^{(i,\delta)}$ defined by the action of $\mathcal{P}^{(i,\delta)}$ on $u = \vec{x} \cdot \vec{\phi}^{(i,\delta)} \in V_h^{(i,\delta)}$ as

$$\mathcal{P}^{(i,\delta)} u = \vec{x} \cdot \mathcal{P}^{(i,\delta)} \vec{\phi}^{(i,\delta)} = \left(\mathbf{P}^{(i,\delta)} \vec{x} \right) \cdot \vec{\phi}.$$

Moreover, define the restriction matrices $\mathbf{R}^{(i,\delta)} := \left(\mathbf{P}^{(i,\delta)} \right)^T$, and the corresponding operator $\mathcal{R}^{(i,\delta)}$.

Let a be a symmetric, continuous, bilinear form on $V \times V$ that is V -elliptic with constant η . For $u = \vec{x} \cdot \vec{\phi}$, $v = \vec{y} \cdot \vec{\phi} \in V_h$,

$$a(u, v) = \vec{x} \cdot a \left(\vec{\phi}, \vec{\phi} \right) \cdot \vec{y} =: \vec{x} \cdot \mathbf{A} \vec{y},$$

where $\mathbf{A}$ is the stiffness matrix. Similarly, we define the subdomain stiffness matrices $\mathbf{A}^{(i,\delta)} := a(\vec{\phi}^{(i,\delta)}, \vec{\phi}^{(i,\delta)})$. Hence, for $u = \vec{x} \cdot \vec{\phi}^{(i,\delta)}, v = \vec{y} \cdot \vec{\phi}^{(i,\delta)} \in V_h^{(i,\delta)}$

$$\begin{aligned} \vec{x} \cdot \mathbf{A}^{(i,\delta)} \cdot \vec{y} &= a\left(\vec{x} \cdot \vec{\phi}^{(i,\delta)}, \vec{y} \cdot \vec{\phi}^{(i,\delta)}\right) \\ &= a\left(\left(\mathbf{P}^{(i,\delta)} \vec{x}\right) \cdot \vec{\phi}, \left(\mathbf{P}^{(i,\delta)} \vec{y}\right) \cdot \vec{\phi}\right) \\ &= \left(\mathbf{P}^{(i,\delta)} \vec{x}\right) \cdot a\left(\vec{\phi}, \vec{\phi}\right) \cdot \left(\mathbf{P}^{(i,\delta)} \vec{y}\right) \\ &= \vec{x} \cdot \mathbf{R}^{(i,\delta)} \mathbf{A} \mathbf{P}^{(i,\delta)} \cdot \vec{y}. \end{aligned}$$

Therefore, $\mathbf{A}^{(i,\delta)} = \mathbf{R}^{(i,\delta)} \mathbf{A} \mathbf{P}^{(i,\delta)}$. We also set $\mathbf{A}^{(i,\delta,\delta+1)} := \mathbf{R}^{(i,\delta)} \mathbf{A} \mathbf{P}^{(i,\delta+1)}$ so that $\mathbf{R}^{(i,\delta)} \mathbf{A} = \mathbf{A}^{(i,\delta,\delta+1)} \mathbf{R}^{(i,\delta+1)}$.

Moreover, let a partition of unity be given by $\theta^{(i,\delta)} \in C(\Omega, [0, 1])$ such that $\sum_i \theta^{(i,\delta)} = 1$ and $\text{supp } \theta^{(i,\delta)} \subset \Omega^{(i,\delta)}$. The operator

$$\mathcal{D}^{(i,\delta)}(u) := \mathcal{I}_h(\theta^{(i,\delta)} u), \quad u \in V_h^{(i,\delta)}$$

has the matrix representation

$$\mathbf{D}^{(i,\delta)} = \text{diag}(\theta^{(i,\delta)}(z_j)).$$

Then the discrete partition of unity property reads as

$$\sum_i \mathbf{P}^{(i,\delta)} \mathbf{D}^{(i,\delta)} \mathbf{R}^{(i,\delta)} = \mathbf{I}_n.$$

A particular choice of such a partition function is $\theta^{(i,-1)}$ such that $\theta^{(i,-1)}(z_j) \in \{0, 1\}$, which leads to an algebraically non-overlapping partition of the unknowns. We write the restriction and prolongation matrices and operators using the superscript $(i, -1)$.

For every subdomain i , we define its set of neighbors by

$$\mathcal{O}_i := \{j \neq i \mid \mathcal{I}^{(i,\delta)} \cap \mathcal{I}^{(j,\delta)} \neq \emptyset\}.$$

Lemma 22.

For $i = 1, \dots, N$

$$\mathbf{P}^{(i,\delta)} \mathbf{A}^{(i,\delta)} \mathbf{D}^{(i,\delta)} = \mathbf{A} \mathbf{P}^{(i,\delta)} \mathbf{D}^{(i,\delta)}, \quad \mathbf{D}^{(i,\delta)} \mathbf{A}^{(i,\delta)} \mathbf{R}^{(i,\delta)} = \mathbf{D}^{(i,\delta)} \mathbf{R}^{(i,\delta)} \mathbf{A}.$$

Proof.

$$\begin{aligned}
D^{(i,\delta)} R^{(i,\delta)} A - D^{(i,\delta)} A^{(i,\delta)} R^{(i,\delta)} &= D^{(i,\delta)} R^{(i,\delta)} a(\vec{\phi}, \phi) - D^{(i,\delta)} a(\phi^{(i)}, \phi^{(i)}) R^{(i,\delta)} \\
&= a(D^{(i,\delta)} \phi^{(i)}, \phi) - a(D^{(i,\delta)} \phi^{(i)}, (P^{(i,\delta)} \phi^{(i)})) \\
&= a(D^{(i,\delta)} \phi^{(i)}, (\phi - P^{(i,\delta)} \phi^{(i)})) \\
&= \left\{ a(\theta(z_j) \phi_j, \delta_{\{k \notin \mathcal{I}^{(i)}\}} \phi_k) \right\}_{j \in \mathcal{I}^{(i)}, k \in \mathcal{I}} \\
&= \left\{ \theta(z_j) \delta_{\{k \notin \mathcal{I}^{(i)}\}} a(\phi_j, \phi_k) \right\}_{j \in \mathcal{I}^{(i)}, k \in \mathcal{I}}
\end{aligned}$$

Now, $a(\phi_j, \phi_k) \neq 0$ only if i and j belong to the same element. For $k \notin \mathcal{I}^{(i)}$, $a(\phi_j, \phi_k) \neq 0$ therefore implies that z_j must be on the boundary of subdomain $\Omega^{(i,\delta)}$. Since θ_i is a continuous function that is zero outside of $\Omega^{(i,\delta)}$, we have $\theta(z_j) = 0$ and therefore that $D^{(i,\delta)} R^{(i,\delta)} A = D^{(i,\delta)} A^{(i,\delta)} R^{(i,\delta)}$. The second equality is the transpose of the first. $\square$

An immediate consequence of the lemma is that

$$\sum_i P^{(i,\delta)} D^{(i,\delta)} A^{(i,\delta)} R^{(i,\delta)} = \sum_i P^{(i,\delta)} D^{(i,\delta)} R^{(i,\delta)} A = A.$$

We define the additive Schwarz preconditioner

$$N_{AS} = \sum_i P^{(i,\delta)} \left(A^{(i,\delta)} \right)^{-1} R^{(i,\delta)},$$

the additive Schwarz preconditioner with harmonic overlap [23]

$$N_{ASH} = \sum_i P^{(i,\delta)} \left(A^{(i,\delta)} \right)^{-1} D^{(i,\delta)} R^{(i,\delta)},$$

and the restricted additive Schwarz preconditioner [24]

$$N_{RAS} = \sum_i P^{(i,\delta)} D^{(i,\delta)} \left(A^{(i,\delta)} \right)^{-1} R^{(i,\delta)}.$$

The corresponding iterative schemes $\mathbb{S}_\bullet$ are given by

$$\vec{x}^{k+1} = \mathbb{S}_{AS}(\vec{b}, \vec{x}^k) := \vec{x}^k + \theta N_{AS}(\vec{b} - A \vec{x}^k),$$

$$\begin{aligned}\vec{x}^{k+1} &= \mathbb{S}_{ASH}(\vec{b}, \vec{x}^k) := \vec{x}^k + \mathbf{N}_{ASH}(\vec{b} - \mathbf{A}\vec{x}^k), \\ \vec{x}^{k+1} &= \mathbb{S}_{RAS}(\vec{b}, \vec{x}^k) := \vec{x}^k + \mathbf{N}_{RAS}(\vec{b} - \mathbf{A}\vec{x}^k)\end{aligned}$$

in global form. The additive Schwarz iteration is only convergent if the damping parameter θ is small enough.

It is typically observed that the number of iterations of the above methods deteriorates with a growing number of subdomains. In order to obtain a scalable solver, the inclusion of a coarse grid is therefore often mandatory. Let $V_H \subset V_h$, and designate the canonical injection by

$$\mathcal{P}^0 : V_H \hookrightarrow V_h.$$

The associated matrix form is called $\mathbf{P}^{(0)}$, and the corresponding restriction is $\mathbf{R}^{(0)} = \left(\mathbf{P}^{(0)}\right)^T$. This permits to define a coarse grid system matrix $\mathbf{A}^{(0)}$ through the Galerkin relationship

$$\mathbf{A}^{(0)} = \mathbf{R}^{(0)} \mathbf{A} \mathbf{P}^{(0)},$$

and a coarse grid preconditioner and correction scheme as

$$\begin{aligned}\mathbf{N}_{CG} &= \mathbf{P}^{(0)} \left(\mathbf{A}^{(0)}\right)^{-1} \mathbf{R}^{(0)}, \\ \vec{x}^{k+1} &= \mathbb{S}_{CG}(\vec{b}, \vec{x}^k) := \vec{x}^k + \mathbf{N}_{CG}(\vec{b} - \mathbf{A}\vec{x}^k).\end{aligned}$$

Two grid methods can then be obtained by combining a domain decomposition iteration with the coarse grid correction, for example as

$$\vec{x}^{k+1} = \left[\mathbb{S}_{CG}(\vec{b}, \cdot) \circ \mathbb{S}_{ASH}(\vec{b}, \cdot) \right] (\vec{x}^k).$$

We will now turn to a detailed description of how faults will affect the convergence properties of one and two level domain decomposition solvers.

8.2 Modeling Blockwise Faults in Iterative Schemes

For the purposes of implementation, and to study the impact of faults, the local form of the above iterations is more helpful. We notice that faults of subdomain solves can not be modeled as before, since, due to the overlap, one unknown can be owned by several subdomains. This means that we cannot model faults by a random diagonal matrix with blockwise dependent diagonal entries.

The additive Schwarz iteration can be written as

$$\begin{aligned}\vec{x}_i^{k+1} &= \vec{x}_i^k + \mathbf{R}^{(i,\delta+1)} \mathbf{N}_{AS} \left(\vec{b} - \mathbf{A} \vec{x}^k \right) \\ &= \vec{x}_i^k + \sum_{j \in \mathcal{O}_i} \mathbf{R}^{(i,\delta+1)} \mathbf{P}^{(j,\delta)} \left(\mathbf{A}^{(j,\delta)} \right)^{-1} \left(\vec{b}_j - \mathbf{A}^{(i,\delta,\delta+1)} \vec{x}_i^k \right),\end{aligned}$$

where $\vec{x}_i^k = \mathbf{R}^{(i,\delta+1)} \vec{x}^k$ and $\vec{b}_i = \mathbf{R}^{(i,\delta)} \vec{b}$ are the node local part of the solution and the right-hand side respectively.

In practice, one subdomain will be assigned to one compute node. We can model the effect of failing compute nodes by Bernoulli distributed random variables

$$\chi^{(i)} = \begin{cases} 1 & \text{with probability } 1 - \varepsilon, \\ 0 & \text{with probability } \varepsilon. \end{cases}$$

Here, the probability of a node failure ε is taken to be small. We rewrite the additive Schwarz scheme introduced above by introducing faults:

$$\vec{x}_i^{k+1} = \vec{x}_i^k + \sum_{j \in \mathcal{O}_i} \mathbf{R}^{(i,\delta+1)} \chi^{(j)} \mathbf{P}^{(j,\delta)} \left(\mathbf{A}^{(j,\delta)} \right)^{-1} \left(\vec{b}_j - \mathbf{A}^{(j,\delta,\delta+1)} \vec{x}_j^k \right).$$

Since all operations up to the communication step are node local, we only need to include one random variable per node to account for any failure on that node.

In order to analyze the convergence behavior, we derive the error propagator $\mathcal{E}_{AS}$ associated with the error equation. The error $\vec{e}^{k+1}$ after $k + 1$ iterations in terms of the error after k iterations by

$$\vec{e}^{k+1} = \vec{x}^{k+1} - \vec{x} = \vec{e}^k + \sum_j \chi^{(j)} \mathbf{P}^{(j,\delta)} \left(\mathbf{A}^{(i,\delta)} \right)^{-1} \mathbf{R}^{(j,\delta)} \mathbf{A} (\vec{x} - \vec{x}^k)$$

$$\begin{aligned}
&= \left(\mathbf{I} - \sum_j \chi^{(j)} \mathbf{P}^{(j,\delta)} \left(\mathbf{A}^{(i,\delta)} \right)^{-1} \mathbf{R}^{(j,\delta)} \mathbf{A} \right) \vec{e}^k \\
&=: \mathcal{E}_{AS} \vec{e}^k.
\end{aligned}$$

The local version of the additive Schwarz method with harmonic overlap reads as

$$\begin{aligned}
\vec{x}_i^{k+1} &= \vec{x}_i^k + \mathbf{R}^{(i,\delta)} \mathbf{N}_{ASH} \left(\vec{b} - \mathbf{A} \vec{x}^k \right) \\
&= \vec{x}_i^k + \sum_{j \in \mathcal{O}_i} \mathbf{R}^{(i,\delta)} \mathbf{P}^{(j,\delta)} \left(\mathbf{A}^{(i,\delta)} \right)^{-1} \mathbf{D}^{(j,\delta)} \left(\mathbf{R}^{(j,\delta)} \vec{b} - \mathbf{R}^{(j,\delta)} \mathbf{A} \vec{x}^k \right) \\
&= \vec{x}_i^k + \sum_{j \in \mathcal{O}_i} \mathbf{R}^{(i,\delta)} \mathbf{P}^{(j,\delta)} \left(\mathbf{A}^{(i,\delta)} \right)^{-1} \mathbf{D}^{(j,\delta)} \left(\vec{b}_j - \mathbf{A}^{(j,\delta)} \vec{x}_j^k \right)
\end{aligned}$$

with $\vec{x}_i^k = \mathbf{R}^{(i,\delta)} \vec{x}^k$, where we used Lemma 22 in the last equality. This shows that again, only one random variable per node is necessary to capture all failures that that node might suffer in one iteration. Similar to the additive Schwarz method, we obtain the error propagator

$$\mathcal{E}_{ASH} = \mathbf{I} - \sum_j \chi^{(j)} \mathbf{P}^{(j,\delta)} \left(\mathbf{A}^{(i,\delta)} \right)^{-1} \mathbf{D}^{(j,\delta)} \mathbf{R}^{(j,\delta)} \mathbf{A}.$$

The reduced additive Schwarz method

$$\begin{aligned}
\vec{x}_i^{k+1} &= \vec{x}_i^k + \mathbf{R}^{(i,\delta)} \mathbf{N}_{RAS} \left(\vec{b} - \mathbf{A} \vec{x}^k \right) \\
&= \vec{x}_i^k + \sum_{j \in \mathcal{O}_i} \mathbf{R}^{(i,\delta)} \mathbf{P}^{(j,\delta)} \mathbf{D}^{(j,\delta)} \left(\mathbf{A}^{(i,\delta)} \right)^{-1} \left(\vec{b}_j - \sum_{l \in \mathcal{O}_j} \mathbf{R}^{(j,\delta)} \mathbf{P}^{(l,\delta)} \mathbf{D}^{(l,\delta)} \mathbf{A}^{(l,\delta)} \vec{x}_l^k \right),
\end{aligned}$$

can be rewritten as

$$\begin{aligned}
\vec{y}_i^0 &= x_i^0, \\
\vec{y}_i^{k+1} &= \vec{y}_i^k + \left(\mathbf{A}^{(i,\delta)} \right)^{-1} \left(\vec{b}_i - \sum_{j \in \mathcal{O}_i} \mathbf{R}^{(i,\delta)} \mathbf{P}^{(j,\delta)} \mathbf{A}^{(j,\delta)} \mathbf{D}^{(j,\delta)} \vec{y}_j^k \right), \\
\vec{x}_i^k &= \sum_{j \in \mathcal{O}_i} \mathbf{R}^{(i,\delta)} \mathbf{P}^{(j,\delta)} \mathbf{D}^{(j,\delta)} \vec{y}_j^k.
\end{aligned}$$

which leads to

$$\mathcal{E}_{RAS} = I - \sum_j \chi^{(j)} \mathbf{P}^{(j,\delta)} \mathbf{D}^{(j,\delta)} \left(\mathbf{A}^{(i,\delta)} \right)^{-1} \mathbf{R}^{(j,\delta)} \mathbf{A}.$$

Finally, the coarse grid correction can be given as

$$\begin{aligned} \vec{x}_i^{k+1} &= \vec{x}_i^k + \mathbf{R}^{(i,\delta)} \mathbf{N}_{CG} \left(\vec{b} - \mathbf{A} \vec{x}^k \right) \\ &= \vec{x}_i^k + \mathbf{R}^{(i,\delta)} \mathbf{P}^{(0)} \left(\mathbf{A}^{(0)} \right)^{-1} \sum_j \mathbf{R}^{(0)} \mathbf{P}^{(j,-1)} \left(\mathbf{R}^{(j,-1)} \vec{b} - \mathbf{R}^{(j,-1)} \mathbf{A} \vec{x}^k \right) \\ &= \vec{x}_i^k + \mathbf{R}^{(i,\delta)} \mathbf{P}^{(0)} \left(\mathbf{A}^{(0)} \right)^{-1} \sum_j \mathbf{R}^{(0)} \mathbf{P}^{(j,-1)} \left(\vec{b}_j - \mathbf{A}^{(j,-1)} \mathbf{R}^{(j,-1)} \mathbf{P}^{(j,\delta)} \vec{x}_j^k \right) \end{aligned}$$

with $\vec{b}_j = \mathbf{R}^{(j,-1)} \vec{b}$. The error propagator of the coarse grid is given by

$$\mathcal{E}_{CG} = I - \mathcal{X}^{(N)} \mathbf{N}_{CG} \sum_j \chi^{(j)} \mathbf{P}^{(j,-1)} \mathbf{A}^{(j,-1)} \mathbf{R}^{(j,-1)}.$$

Here, two sources of faults are possible: faults can occur during the computation of the residual, or during the coarse grid solve and update. The random variables $\chi^{(j)}$ model the former, whereas the diagonal fault matrix $\mathcal{X}^{(N)}$ encodes the fault structure of the coarse grid solve. Since we chose an algebraically non-overlapping partition, we can write

$$\mathcal{E}_{CG} = I - \mathcal{X}^{(N)} \mathbf{N}_{CG} \mathcal{X}^{(\rho)} \mathbf{A},$$

where the diagonal faults matrix $\mathcal{X}^{(\rho)}$ has block structure:

$$\begin{aligned} \mathcal{X}_{ii}^{(\rho)} &= \chi^{(j)} && \text{if } \theta^{(j,-1)}(z_i) = 1, \\ \mathcal{X}_{ik}^{(\rho)} &= 0 && \text{if } i \neq k. \end{aligned}$$

The corresponding error propagators for the classical fault free methods are obtained by replacing the random variables with identity, and are designated by $\mathbf{E}_{AS}$, $\mathbf{E}_{ASH}$, $\mathbf{E}_{RAS}$ and $\mathbf{E}_{CG}$ respectively.

8.3 Convergence Analysis for Domain Decomposition Methods

Our first result concerns the convergence behavior of the single level domain decomposition methods introduced above.

A standard assumption needed in the convergence analysis of fault-free domain decomposition methods concerns the interaction of the local solution spaces $V^{(i,\delta)}$:

Assumption 6 (Strengthened Cauchy-Schwarz inequality): Define $0 \leq \epsilon_{ij} \leq 1$ to be the smallest constants such that

$$|a(u, v)| \leq \epsilon_{ij} a(u, u)^{1/2} a(v, v)^{1/2} \quad \forall u \in V^{(i,\delta)}, v \in V^{(j,\delta)}.$$

If N_c denotes the smallest number of colors necessary to color the subdomains, then

$$\rho(\epsilon) \leq N_c,$$

and N_c is independent of h and N .

Theorem 23.

Let $\mathcal{E}_{DD} \in \{\mathcal{E}_{AS}, \mathcal{E}_{ASH}, \mathcal{E}_{RAS}\}$, and let $\mathbf{E}_{DD}$ the iteration matrix of its fault-free equivalent. Provided Assumption 6 holds, the convergence rate of the fault-prone iterations is bounded as

$$\varrho(\mathcal{E}_{DD}) \leq \|\mathbf{E}_{DD}\|_A + C\epsilon,$$

where $\|\bullet\|_A$ is the energy norm and C is independent of ϵ , h and N .

Proof. For easy of notation, we define the extended matrices

$$\begin{aligned} \overline{\mathbf{A}} &:= \text{blockdiag} \left(\mathbf{A}^{(i,\delta)} \right), & \overline{\mathbf{D}} &:= \text{blockdiag} \left(\mathbf{D}^{(i,\delta)} \right), \\ \overline{\mathbf{R}} &:= \text{vec} \left(\mathbf{R}^{(i,\delta)} \right), & \overline{\mathbf{P}} &:= \text{vec} \left(\mathbf{P}^{(i,\delta)} \right), \\ \overline{\mathbf{X}} &:= \text{blockdiag} \left(\chi^{(i)} \mathbf{I}_{n^{(i,\delta)}} \right), \end{aligned}$$

and write

$$\begin{aligned}
\mathbf{E}_{AS} &= \mathbf{I} - \overline{\mathbf{P}\mathbf{A}}^{-1}\overline{\mathbf{R}\mathbf{A}}, & \mathbf{E}_{AS} &= \mathbf{I} - \overline{\mathbf{P}\mathbf{X}\mathbf{A}}^{-1}\overline{\mathbf{R}\mathbf{A}}, \\
\mathbf{E}_{ASH} &= \mathbf{I} - \overline{\mathbf{P}\mathbf{A}}^{-1}\overline{\mathbf{D}\mathbf{R}\mathbf{A}}, & \mathbf{E}_{ASH} &= \mathbf{I} - \overline{\mathbf{P}\mathbf{X}\mathbf{A}}^{-1}\overline{\mathbf{D}\mathbf{R}\mathbf{A}}, \\
\mathbf{E}_{RAS} &= \mathbf{I} - \overline{\mathbf{P}\mathbf{D}\mathbf{A}}^{-1}\overline{\mathbf{R}\mathbf{A}}, & \mathbf{E}_{RAS} &= \mathbf{I} - \overline{\mathbf{P}\mathbf{X}\mathbf{D}\mathbf{A}}^{-1}\overline{\mathbf{R}\mathbf{A}}.
\end{aligned}$$

Using the fact that $\overline{\mathbf{A}}^{1/2}$ and $\overline{\mathbf{X}}$ commute, we find

$$\begin{aligned}
\|\mathbb{V}[\mathbf{E}_{AS}]\|_A &= \left\| \overline{\mathbf{P}}^{\otimes 2} \mathbb{V}[\overline{\mathbf{X}}] \left(\overline{\mathbf{A}}^{-1} \overline{\mathbf{R}\mathbf{A}} \right)^{\otimes 2} \right\|_A \\
&= \left\| \left(\overline{\mathbf{A}}^{1/2} \overline{\mathbf{P}} \right)^{\otimes 2} \mathbb{V}[\overline{\mathbf{X}}] \left(\overline{\mathbf{A}}^{-1} \overline{\mathbf{R}\mathbf{A}}^{1/2} \right)^{\otimes 2} \right\|_2 \\
&= \left\| \left(\overline{\mathbf{A}}^{1/2} \overline{\mathbf{P}\mathbf{A}}^{-1/2} \right)^{\otimes 2} \mathbb{V}[\overline{\mathbf{X}}] \left(\overline{\mathbf{A}}^{-1/2} \overline{\mathbf{R}\mathbf{A}}^{1/2} \right)^{\otimes 2} \right\|_2 \\
&\leq \left\| \overline{\mathbf{A}}^{1/2} \overline{\mathbf{P}\mathbf{A}}^{-1/2} \right\|_2^2 \|\mathbb{V}[\overline{\mathbf{X}}]\|_2 \left\| \overline{\mathbf{A}}^{-1/2} \overline{\mathbf{R}\mathbf{A}}^{1/2} \right\|_2^2 \\
&\leq C\varepsilon\rho(\mathbf{N}_{AS}\mathbf{A})^2.
\end{aligned}$$

Similarly, we find

$$\begin{aligned}
\|\mathbb{V}[\mathbf{E}_{ASH}]\|_A &= \left\| \overline{\mathbf{P}}^{\otimes 2} \mathbb{V}[\overline{\mathbf{X}}] \left(\overline{\mathbf{A}}^{-1} \overline{\mathbf{D}\mathbf{R}\mathbf{A}} \right)^{\otimes 2} \right\|_A \\
&= \left\| \left(\overline{\mathbf{P}\mathbf{A}}^{-1/2} \right)^{\otimes 2} \mathbb{V}[\overline{\mathbf{X}}] \left(\overline{\mathbf{A}}^{-1/2} \overline{\mathbf{D}\mathbf{R}\mathbf{A}} \right)^{\otimes 2} \right\|_A \\
&= \left\| \left(\mathbf{A}^{1/2} \overline{\mathbf{P}\mathbf{A}}^{-1/2} \right)^{\otimes 2} \mathbb{V}[\overline{\mathbf{X}}] \left(\overline{\mathbf{A}}^{-1/2} \overline{\mathbf{D}\mathbf{R}\mathbf{A}}^{1/2} \right)^{\otimes 2} \right\|_2 \\
&\leq \left\| \mathbf{A}^{1/2} \overline{\mathbf{P}\mathbf{A}}^{-1/2} \right\|_2^2 \|\mathbb{V}[\overline{\mathbf{X}}]\|_2 \left\| \overline{\mathbf{A}}^{-1/2} \overline{\mathbf{D}\mathbf{R}\mathbf{A}}^{1/2} \right\|_2^2 \\
&= \rho\left(\overline{\mathbf{P}\mathbf{A}}^{-1} \overline{\mathbf{R}\mathbf{A}}\right) \|\mathbb{V}[\overline{\mathbf{X}}]\|_2 \rho\left(\mathbf{A}^{1/2} \overline{\mathbf{P}\mathbf{D}\mathbf{A}}^{-1} \overline{\mathbf{D}\mathbf{R}\mathbf{A}}^{1/2}\right) \\
&= \rho\left(\overline{\mathbf{P}\mathbf{A}}^{-1} \overline{\mathbf{R}\mathbf{A}}\right) \|\mathbb{V}[\overline{\mathbf{X}}]\|_2 \rho\left(\overline{\mathbf{P}\mathbf{D}\mathbf{A}}^{-1} \overline{\mathbf{D}\mathbf{R}\mathbf{A}}\right) \\
&\leq C\varepsilon\rho(\mathbf{N}_{AS}\mathbf{A})\rho(\mathbf{N}_{RASH}\mathbf{A})
\end{aligned}$$

where we have set

$$\mathbf{N}_{RASH} = \overline{\mathbf{P}\mathbf{D}\mathbf{A}}^{-1} \overline{\mathbf{D}\mathbf{R}} = \sum_i \mathbf{P}^{(i,\delta)} \mathbf{D}^{(i,\delta)} \left(\mathbf{A}^{(i,\delta)} \right)^{-1} \mathbf{D}^{(i,\delta)} \mathbf{R}^{(i,\delta)}.$$

It can be shown that the same holds for the reduced additive Schwarz method. Therefore we have shown that

$$\begin{aligned}\varrho(\mathcal{E}_{AS}) &\leq \sqrt{(\|\mathbf{E}_{AS}\|_A + \varepsilon \|\mathbf{N}_{AS}\mathbf{A}\|_A)^2 + C\varepsilon\rho(\mathbf{N}_{AS}\mathbf{A})^2}, \\ \varrho(\mathcal{E}_{ASH}) &\leq \sqrt{(\|\mathbf{E}_{ASH}\|_A + \varepsilon \|\mathbf{N}_{ASH}\mathbf{A}\|_A)^2 + C\varepsilon\rho(\mathbf{N}_{AS}\mathbf{A})\rho(\mathbf{N}_{RASH}\mathbf{A})}, \\ \varrho(\mathcal{E}_{RAS}) &\leq \sqrt{(\|\mathbf{E}_{RAS}\|_A + \varepsilon \|\mathbf{N}_{RAS}\mathbf{A}\|_A)^2 + C\varepsilon\rho(\mathbf{N}_{AS}\mathbf{A})\rho(\mathbf{N}_{RASH}\mathbf{A})}.\end{aligned}$$

We conclude using Lemmas 24 and 25 below. $\square$

Lemma 24.

The partition of unity $\theta^{(i,\delta)}$ can be chosen such that $\|\nabla\theta^{(i,\delta)}\|_{L^\infty} \leq C/(\delta h)$. Then, for all $u \in V_h^{(i,\delta)}$,

$$a(\mathcal{D}^{(i,\delta)}u, \mathcal{D}^{(i,\delta)}u) \leq C_D a(u, u)$$

for some constant C_D .

Proof. Set $B^{(i,\delta)} := \Omega^{(i,\delta)} \cap \{\theta^{(i,\delta)} < 1\}$. For $u \in V_h^{(i,\delta)}$, we have that

$$\begin{aligned}a(\mathcal{D}^{(i,\delta)}u, \mathcal{D}^{(i,\delta)}u) &\leq C |\mathcal{D}^{(i,\delta)}u|_{H^1(\Omega^{(i,\delta)})}^2 \\ &= C \left(|u|_{H^1(\Omega^{(i,\delta)} \setminus B^{(i,\delta)})}^2 + |\mathcal{I}_h[\theta^{(i,\delta)}u]|_{H^1(B^{(i,\delta)})}^2 \right).\end{aligned}$$

The second term is estimated as

$$\begin{aligned}|\mathcal{I}_h[\theta^{(i,\delta)}u]|_{H^1(B^{(i,\delta)})} &\leq |\theta^{(i,\delta)}u|_{H^1(B^{(i,\delta)})} + |\theta^{(i,\delta)}u - \mathcal{I}_h[\theta^{(i,\delta)}u]|_{H^1(B^{(i,\delta)})} \\ &\leq C |\theta^{(i,\delta)}u|_{H^1(B^{(i,\delta)})}.\end{aligned}$$

Now

$$\begin{aligned}|\theta^{(i,\delta)}u|_{H^1(B^{(i,\delta)})}^2 &\leq 2 \int_{B^{(i,\delta)}} (\theta^{(i,\delta)})^2 |\nabla u|^2 + 2 \int_{B^{(i,\delta)}} u^2 |\nabla\theta^{(i,\delta)}|^2 \\ &\leq 2 |u|_{H^1(B^{(i,\delta)})}^2 + 2 \|\nabla\theta^{(i,\delta)}\|_{L^\infty}^2 \|u\|_{L^2(B^{(i,\delta)})}^2.\end{aligned}\tag{8.1}$$

Application of a Poincare inequality then leads to

$$|\theta^{(i,\delta)}u|_{H^1(B^{(i,\delta)})}^2 \leq 2 |u|_{H^1(B^{(i,\delta)})}^2 + 2C \|\nabla\theta^{(i,\delta)}\|_{L^\infty}^2 (\delta h)^2 |u|_{H^1(B^{(i,\delta)})}^2$$

$$\leq 2 |u|_{H^1(B^{(i,\delta)})}^2 + 2C |u|_{H^1(B^{(i,\delta)})}^2.$$

Therefore, we have shown

$$\begin{aligned} a(\mathcal{D}^{(i,\delta)}u, \mathcal{D}^{(i,\delta)}u) &\leq C |u|_{H^1(\Omega^{(i,\delta)})}^2 \\ &\leq Ca(u, u). \end{aligned}$$

□

Lemma 25.

$$\begin{aligned} \rho(\mathbf{N}_{AS}\mathbf{A}) &= \|\mathbf{N}_{AS}\mathbf{A}\|_A \leq N_c, \\ \rho(\mathbf{N}_{ASH}\mathbf{A}) &\leq \|\mathbf{N}_{ASH}\mathbf{A}\|_A \leq N_c C_D^{1/2}, \\ \rho(\mathbf{N}_{RAS}\mathbf{A}) &\leq \|\mathbf{N}_{RAS}\mathbf{A}\|_A \leq N_c C_D^{1/2}, \\ \rho(\mathbf{N}_{RASH}\mathbf{A}) &\leq \|\mathbf{N}_{RASH}\mathbf{A}\|_A \leq N_c C_D. \end{aligned}$$

Proof. The proof of the first inequality is classical and makes use of the operator form $\mathcal{T}_{AS}$ of $\mathbf{N}_{AS}\mathbf{A}$. This operator as well as the ones for $\mathbf{N}_{ASH}\mathbf{A}$, $\mathbf{N}_{RAS}\mathbf{A}$ and $\mathbf{N}_{RASH}\mathbf{A}$ are given by

$$\begin{aligned} \mathcal{T}_{AS} &= \sum_i \mathcal{T}^{(i,\delta)}, \\ \mathcal{T}_{ASH} &= \sum_i \tilde{\mathcal{T}}^{(i,\delta)}, \\ \mathcal{T}_{RAS} &= \sum_i \mathcal{D}^{(i,\delta)} \mathcal{T}^{(i,\delta)}, \\ \mathcal{T}_{RASH} &= \sum_i \mathcal{D}^{(i,\delta)} \tilde{\mathcal{T}}^{(i,\delta)}, \end{aligned}$$

where the subdomain solution operators $\mathcal{T}^{(i,\delta)}$ and $\tilde{\mathcal{T}}^{(i,\delta)}$ are defined by

$$\begin{aligned} a(\mathcal{T}^{(i,\delta)}u, v) &= a(u, v) & \forall u \in V_h, v \in V_h^{(i,\delta)}, \\ a(\tilde{\mathcal{T}}^{(i,\delta)}u, v) &= a(u, \mathcal{D}^{(i,\delta)}v) & \forall u \in V_h, v \in V_h^{(i,\delta)}. \end{aligned}$$

We then have

$$\begin{aligned}
\rho(\mathbf{N}_\bullet \mathbf{A})^2 &\leq \|\mathbf{N}_\bullet \mathbf{A}\|_A^2 \\
&= \sup_{\vec{x}} \frac{(\mathbf{N}_\bullet \mathbf{A} \vec{x}) \cdot (\mathbf{A} \mathbf{N}_\bullet \mathbf{A} \vec{x})}{\vec{x} \cdot \mathbf{A} \vec{x}} \\
&= \sup_{u \in V_h} \frac{a(\mathcal{T}_\bullet u, \mathcal{T}_\bullet u)}{a(u, u)}.
\end{aligned}$$

For $\mathcal{T}_\bullet = \sum_i \mathcal{T}_\bullet^{(i)}$, we have by the strengthened Cauchy-Schwarz inequality that

$$\begin{aligned}
a(\mathcal{T}_\bullet u, \mathcal{T}_\bullet u) &= \sum_{i,j} a(\mathcal{T}_\bullet^{(i)} u, \mathcal{T}_\bullet^{(j)} u) \\
&\leq \sum_{i,j} \epsilon_{ij} a(\mathcal{T}_\bullet^{(i)} u, \mathcal{T}_\bullet^{(i)} u)^{1/2} a(\mathcal{T}_\bullet^{(j)} u, \mathcal{T}_\bullet^{(j)} u)^{1/2} \\
&\leq \rho(\epsilon) \sum_i a(\mathcal{T}_\bullet^{(i)} u, \mathcal{T}_\bullet^{(i)} u).
\end{aligned}$$

For $\mathcal{T}_\bullet = \mathcal{T}_{AS}$ and $\mathcal{T}_\bullet^{(i)} = \mathcal{T}^{(i,\delta)}$, we can use that $\mathcal{T}^{(i,\delta)}$ are idempotent and self-adjoint with respect to $a(\cdot, \cdot)$ and obtain

$$\begin{aligned}
a(\mathcal{T}_{AS} u, \mathcal{T}_{AS} u) &\leq \rho(\epsilon) \sum_i a(u, \mathcal{T}^{(i,\delta)} u) \\
&= \rho(\epsilon) a(u, \mathcal{T}_{AS} u) \\
&\leq \rho(\epsilon) a(u, u)^{1/2} a(\mathcal{T}_{AS} u, \mathcal{T}_{AS} u)^{1/2}.
\end{aligned}$$

Therefore $\rho(\mathbf{N}_{AS} \mathbf{A}) \leq \rho(\epsilon)$.

Similarly, for $\mathcal{T}_\bullet = \mathcal{T}_{RAS}$ and $\mathcal{T}_\bullet^{(i)} = \mathcal{D}^{(i,\delta)} \mathcal{T}^{(i,\delta)}$, we have by Lemma 24

$$\begin{aligned}
a(\mathcal{T}_{RAS} u, \mathcal{T}_{RAS} u) &\leq \rho(\epsilon) \sum_i a(\mathcal{D}^{(i,\delta)} \mathcal{T}^{(i,\delta)} u, \mathcal{D}^{(i,\delta)} \mathcal{T}^{(i,\delta)} u) \\
&\leq \rho(\epsilon) C_D \sum_i a(\mathcal{T}^{(i,\delta)} u, \mathcal{T}^{(i,\delta)} u) \\
&\leq \rho(\epsilon) C_D a(u, u)^{1/2} a(\mathcal{T}_{AS} u, \mathcal{T}_{AS} u)^{1/2} \\
&\leq \rho(\epsilon)^2 C_D a(u, u)
\end{aligned}$$

which permits to conclude that $\rho(\mathbf{N}_{RAS} \mathbf{A}) \leq \rho(\epsilon) C_D^{1/2}$.

For $\mathcal{T}_\bullet = \mathcal{T}_{RASH}$ and $\mathcal{T}_\bullet^{(i)} = \mathcal{D}^{(i,\delta)} \tilde{\mathcal{T}}^{(i,\delta)}$, we find

$$\begin{aligned}
a(\mathcal{T}_{RASH}u, \mathcal{T}_{RASH}u) &\leq \rho(\epsilon) \sum_i a\left(\mathcal{D}^{(i,\delta)} \tilde{\mathcal{T}}^{(i,\delta)}u, \mathcal{D}^{(i,\delta)} \tilde{\mathcal{T}}^{(i,\delta)}u\right) \\
&\leq \rho(\epsilon) C_D \sum_i a\left(\tilde{\mathcal{T}}^{(i,\delta)}u, \tilde{\mathcal{T}}^{(i,\delta)}u\right) \\
&= \rho(\epsilon) C_D \sum_i a\left(u, \mathcal{D}^{(i,\delta)} \tilde{\mathcal{T}}^{(i,\delta)}u\right) \\
&= \rho(\epsilon) C_D a(u, \mathcal{T}_{RASH}u) \\
&\leq \rho(\epsilon) C_D a(u, u)^{1/2} a(\mathcal{T}_{RASH}u, \mathcal{T}_{RASH}u)^{1/2}
\end{aligned}$$

and hence $\rho(\mathbf{N}_{RASH}\mathbf{A}) \leq \rho(\epsilon) C_D$.

Finally, for $\mathcal{T}_\bullet = \mathcal{T}_{ASH}$ and $\mathcal{T}_\bullet^{(i)} = \tilde{\mathcal{T}}^{(i,\delta)}$,

$$\begin{aligned}
a(\mathcal{T}_{ASH}u, \mathcal{T}_{ASH}u) &\leq \rho(\epsilon) \sum_i a\left(\tilde{\mathcal{T}}^{(i,\delta)}u, \tilde{\mathcal{T}}^{(i,\delta)}u\right) \\
&\leq \rho(\epsilon) \sum_i a\left(u, \mathcal{D}^{(i,\delta)} \tilde{\mathcal{T}}^{(i,\delta)}u\right) \\
&= \rho(\epsilon) a(u, \mathcal{T}_{RASH}u) \\
&\leq \rho(\epsilon) a(u, u)^{1/2} a(\mathcal{T}_{TASH}u, \mathcal{T}_{TASH}u)^{1/2} \\
&\leq \rho(\epsilon)^2 C_D a(u, u),
\end{aligned}$$

where the previous inequality was used. Therefore $\rho(\mathbf{N}_{RASH}\mathbf{A}) \leq \rho(\epsilon) C_D$. $\square$

In order to verify the above result, we consider the following test case:

$$\text{Find } u \in V_h : \int_{\Omega} \nabla u \cdot \nabla v = \int_{\Omega} f v \quad \forall v \in V_h,$$

where $\Omega = [0, 1]^2$, V_h the space of piecewise linears on a uniform triangular mesh and $f \in L^2(\Omega)$. The resulting linear system is solved using the additive Schwarz method with harmonic overlap on $N = 16$ subdomains for varying mesh sizes h and fault rates ϵ . We iterate for 1000 iterations, as to get a good estimate of the asymptotic convergence rate, and renormalize the residual between iterations. The resulting approximation of $\varrho(\mathcal{E}_{ASH})$ is shown in Figure 8.1. It can be observed that

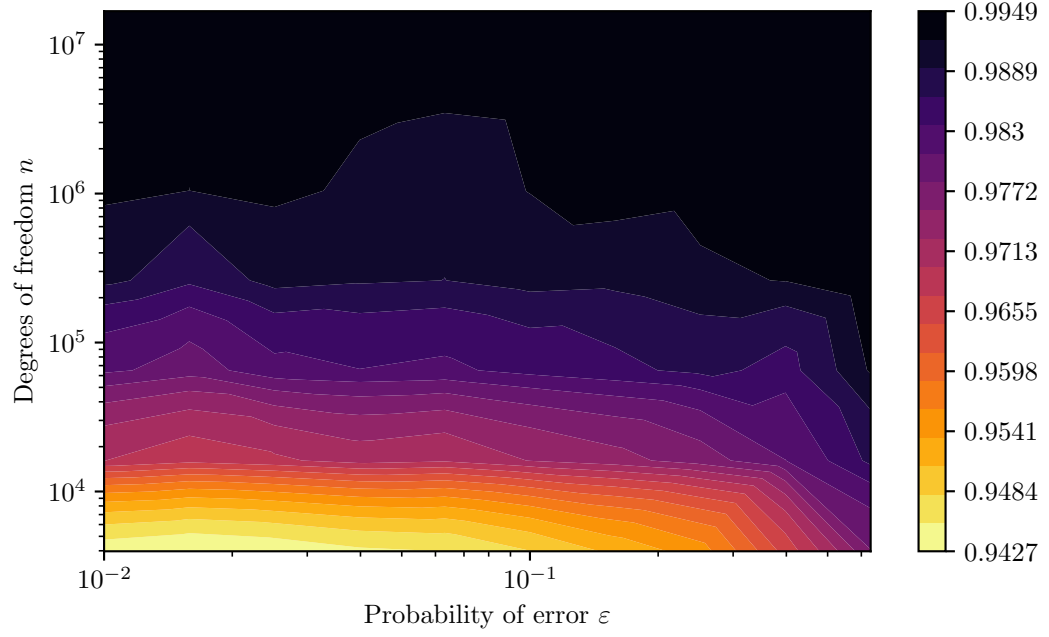


Figure 8.1: Asymptotic convergence rate $\varrho(\mathcal{E}_{ASH})$ for $N = 16$ subdomains.

the convergence rate varies mainly in the dependence on the number of degrees of freedom n , but is largely unaffected by the fault rate ε .

If the number of subdomains N gets scaled proportionally to the number of degrees of freedom n , similar behavior is observed. (See Figure 8.2.) In both cases, the method stays convergent even for considerably large fault rates ε , and essentially mirror their fault free equivalent. This shows that the estimate in Theorem 23 is sharp.

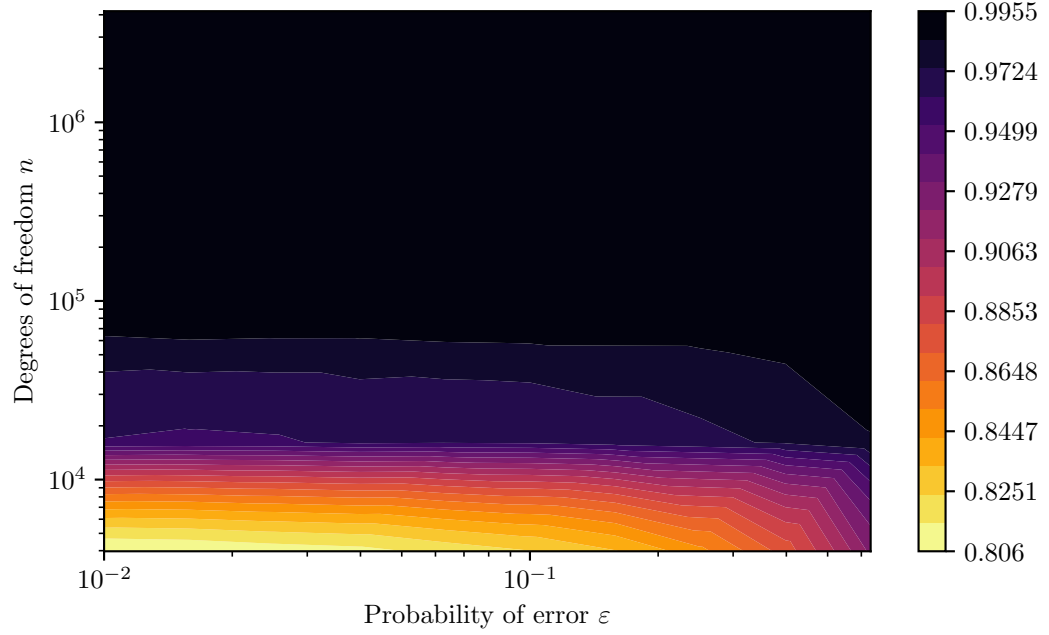


Figure 8.2: Asymptotic convergence rate $\rho(\mathcal{E}_{ASH})$. The number of subdomains N is proportional to the number of degrees of freedom n .

8.4 Two-Level Methods

We now turn to the case when a coarse grid is incorporated into the solver. It has already been established in Chapters 5 and 6 for multigrid solvers that in order to obtain a fault-resilient method, protection of the prolongation from the coarse grid to the fine grid is mandatory. For the same test problem as above, we numerically verify that this also the case for two level domain decomposition solvers. (See Figures 8.3 and 8.4.) Obviously, the single node case does not require protection of the prolongation, since a fault in the coarse solve simply leads to a skipped coarse grid correction. If the number of subdomains is fixed to $N = 16$, divergence can be observed that for large values of ε . The dependence on the number of degrees of freedom seems to be very mild. For the more realistic case of $N \sim n$ seen in Figure 8.4, we observe however that the problem size for which the two level method is convergent scales with $\varepsilon n^{1/2}$. This behavior mirrors what was derived in

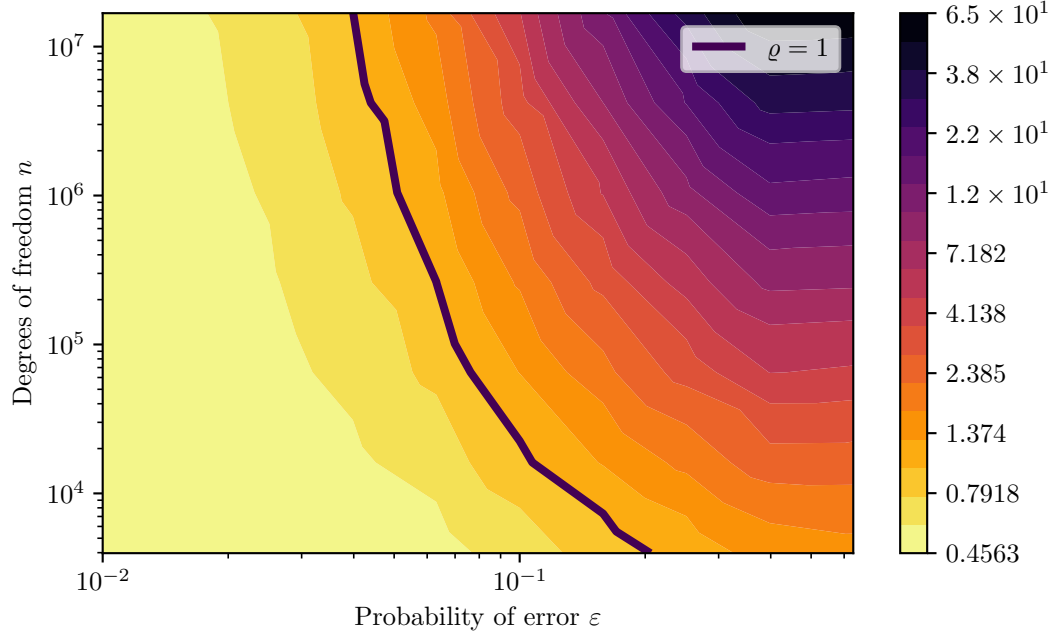


Figure 8.3: Asymptotic convergence rate $\rho(\mathcal{E}_{ASH}\mathcal{E}_{CG})$ without protected prolongation for $N = 16$ subdomains.

Chapter 6 for the multigrid case.

The next result concerns the variance of the coarse grid solver if the prolongation is protected.

Lemma 26.

Take η_{ij} to be the smallest constants such that

$$|\vec{x} \cdot \mathbf{A}^{-1} \vec{y}| \leq \eta_{ij} |\vec{x} \cdot \mathbf{A}^{-1} \vec{x}|^{1/2} |\vec{y} \cdot \mathbf{A}^{-1} \vec{y}|^{1/2}$$

for all $\vec{x} \in \text{range } \mathbf{P}^{(j,-1)}$, $\vec{y} \in \text{range } \mathbf{P}^{(i,-1)}$. Then, if prolongation is protected, we can bound

$$\|\mathbb{V}[\mathcal{E}_{CG}]\|_A \leq C\epsilon h^{-1} \sqrt{\max_i \sum_j \eta_{ij}^2} \leq C\epsilon h^{-1} \sqrt{N}, \quad (8.2)$$

where the tensor energy norm is given by $\|\bullet\|_A := \left\| \left(\mathbf{A}^{1/2} \right)^{\otimes 2} \bullet \left(\mathbf{A}^{-1/2} \right)^{\otimes 2} \right\|_2$.

The proof of Lemma 26 requires several definitions and estimates that concern the block structure of the variance of the fault matrix.

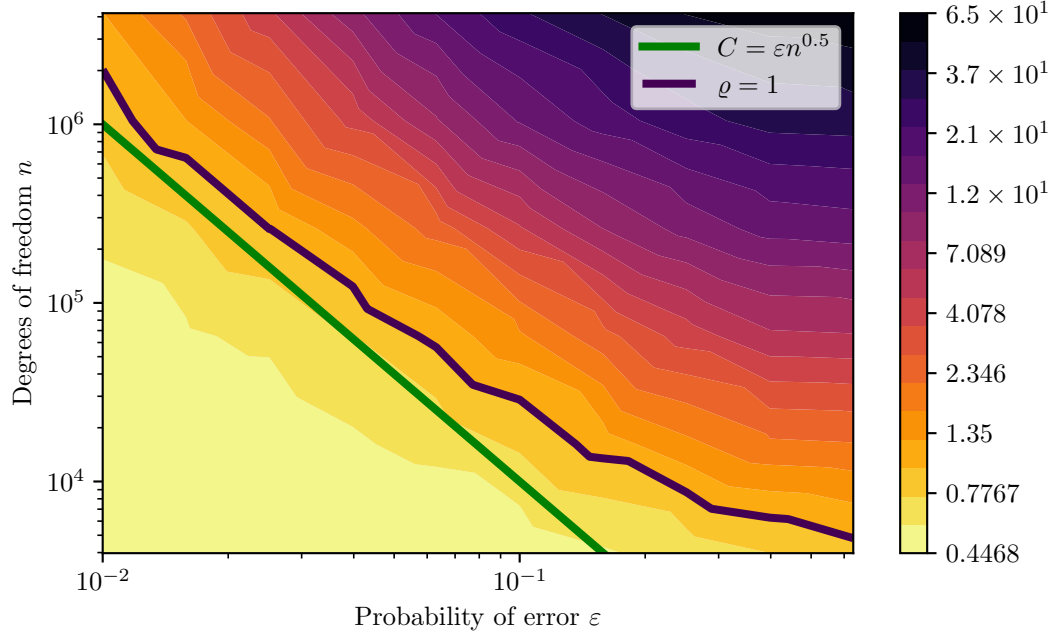


Figure 8.4: Asymptotic convergence rate $\rho(\mathcal{E}_{ASH}\mathcal{E}_{CG})$ without protected prolongation for $N \sim n$.

The variance of $\mathcal{X}^{(\rho)}$ is a diagonal matrix as well, and is given by

$$\mathbb{V}[\mathcal{X}^{(\rho)}]_{in+j,in+j} = \begin{cases} 1 & \text{if } \theta^{(i,-1)}(x_j) = 1 \\ 0 & \text{else} \end{cases}.$$

After reordering, it is a block diagonal matrix with identity in the blocks corresponding to interactions within a node, and zero blocks elsewhere. Before we introduce the notation for the treatment of block matrices, we show the

Example 4.

Consider the case of two subdomains, where the degrees of freedom of the first subdomain are ordered first, and then the unknowns of the second partition. We then have

for the variance of $\mathcal{X}^{(\rho)}$ the form

$$\mathbb{V}[\mathcal{X}^{(\rho)}] = \varepsilon(1 - \varepsilon) \begin{pmatrix} \mathbf{I} & & & \\ & \mathbf{0} & & \\ & & \mathbf{0} & \\ & & & \mathbf{I} \end{pmatrix}.$$

If a matrix $\mathbf{Z} = \begin{pmatrix} \mathbf{Z}_{11} & \mathbf{Z}_{12} \\ \mathbf{Z}_{21} & \mathbf{Z}_{22} \end{pmatrix} \in \mathbb{R}^{n \times n}$ is block partitioned with respect to the sub-domains, we have

$$\begin{aligned} \rho \left(\mathbb{V}[\mathcal{X}^{(\rho)}] \mathbf{Z}^{\otimes 2} \mathbb{V}[\mathcal{X}^{(\rho)}] \right) &= \varepsilon^2(1 - \varepsilon)^2 \rho \left[\begin{pmatrix} \mathbf{Z}_{11}^{\otimes 2} & & \mathbf{Z}_{12}^{\otimes 2} \\ & \mathbf{0} & \\ & & \mathbf{0} \\ \mathbf{Z}_{21}^{\otimes 2} & & \mathbf{Z}_{22}^{\otimes 2} \end{pmatrix} \right] \\ &= \varepsilon^2(1 - \varepsilon)^2 \rho \left[\begin{pmatrix} \mathbf{Z}_{11}^{\otimes 2} & \mathbf{Z}_{12}^{\otimes 2} \\ \mathbf{Z}_{21}^{\otimes 2} & \mathbf{Z}_{22}^{\otimes 2} \end{pmatrix} \right]. \end{aligned}$$

The last block matrix is called the Khatri-Rao square power of $\mathbf{Z}$, and we are interested in estimating its spectral radius.

Definition 27 (Blocks of matrices).

Given ordered vectors of indices $1 \leq \alpha_1 < \alpha_2 < \dots < \alpha_{\ell(\alpha)} \leq m$, $1 \leq \beta_1 < \beta_2 < \dots < \beta_{\ell(\beta)} \leq n$, and a matrix $\mathbf{A} \in \mathbb{R}^{m \times n}$, we designate the sub-matrix corresponding to α and β by

$$\mathbf{A}[\alpha, \beta] = \{ \mathbf{A}_{\alpha_i, \beta_j} \}_{i,j}.$$

If $M = N$ and $\alpha = \beta$, we write

$$\mathbf{A}[\alpha] := \mathbf{A}[\alpha, \alpha].$$

We write the length of the index vector α as $\ell(\alpha)$, so that $\mathbf{A}[\alpha, \beta]$ has size $\ell(\alpha) \times \ell(\beta)$.

Definition 28 (Partitioned matrix).

Let $\mathcal{A}$ and $\mathcal{B}$ be sets of index vectors as above, with $N_{\mathcal{A}} = |\mathcal{A}|$ and $N_{\mathcal{B}} = |\mathcal{B}|$ the respective number of blocks, such that $\bigcup_{\alpha \in \mathcal{A}} \alpha = \{1, \dots, m\}$ with $\alpha \cap \gamma = \emptyset$ for $\alpha \neq \gamma$, and $\bigcup_{\beta \in \mathcal{B}} \beta = \{1, \dots, n\}$ with $\beta \cap \eta = \emptyset$ for $\beta \neq \eta$. We call $\mathbf{A} \in \mathbb{R}^{m \times n}$ partitioned by $\mathcal{A} \times \mathcal{B}$ into the blocks $\mathbf{A}[\alpha, \beta]$, $\alpha \in \mathcal{A}$, $\beta \in \mathcal{B}$. If we prescribe an ordering of $\mathcal{A} = \{\alpha_i\}_{i=1}^{N_{\mathcal{A}}}$ and of $\mathcal{B} = \{\beta_j\}_{j=1}^{N_{\mathcal{B}}}$, we can define the block-ordered permutation $\overline{\mathbf{A}}$ of $\mathbf{A}$:

$$\overline{\mathbf{A}}_{ij} = \mathbf{A}[\alpha_i, \beta_j].$$

Definition 29 (Khatri-Rao product).

Let $\mathbf{A}, \mathbf{B} \in \mathbb{R}^{m \times n}$ be compatibly partitioned by $\mathcal{A} \times \mathcal{B}$. The Khatri-Rao product $\mathbf{A} * \mathbf{B}$ is defined block-wise as

$$(\mathbf{A} * \mathbf{B})_{ij} = \mathbf{A}[\alpha_i, \beta_j] \otimes \mathbf{B}[\alpha_i, \beta_j], \quad 1 \leq i \leq N_{\mathcal{A}}, \quad 1 \leq j \leq N_{\mathcal{B}}.$$

Then $\mathbf{A} * \mathbf{B} \in \mathbb{R}^{\sum_{i=1}^{N_{\mathcal{A}}} \ell(\alpha_i)^2 \times \sum_{j=1}^{N_{\mathcal{B}}} \ell(\beta_j)^2}$. We define the k -th Khatri-Rao power of $\mathbf{A}$ by

$$\mathbf{A}^{*k} := \underbrace{\mathbf{A} * \mathbf{A} * \dots * \mathbf{A}}_{k \text{ times}}.$$

Lemma 30 ([74]).

Let $\mathbf{A}, \mathbf{B}$ be compatibly partitioned positive (semi-)definite matrices. Then $\mathbf{A} * \mathbf{B}$ is positive (semi-)definite.

Lemma 31.

Let $\mathbf{A}, \mathbf{B} \in \mathbb{R}^{n \times m}$ be compatibly partitioned matrices. Then

$$\|\mathbf{A} * \mathbf{B}\|_2 \leq \|(\mathbf{A}^T \mathbf{A}) * \mathbf{I}_m\|_2^{\frac{1}{2}} \|\mathbf{I}_n * (\mathbf{B} \mathbf{B}^T)\|_2^{\frac{1}{2}}.$$

Proof. The matrices $\begin{pmatrix} \mathbf{A} & \mathbf{I}_n \end{pmatrix}^T \begin{pmatrix} \mathbf{A} & \mathbf{I}_n \end{pmatrix}$ and $\begin{pmatrix} \mathbf{I}_m & \mathbf{B}^T \end{pmatrix}^T \begin{pmatrix} \mathbf{I}_m & \mathbf{B}^T \end{pmatrix}$ are trivially positive semi-definite. Using lemma 30 their Khatri-Rao product

$$\begin{pmatrix} (\mathbf{A}^T \mathbf{A}) * \mathbf{I}_m & \mathbf{A}^T * \mathbf{B}^T \\ \mathbf{A} * \mathbf{B} & \mathbf{I}_n * (\mathbf{B} \mathbf{B}^T) \end{pmatrix}$$

is positive semi-definite as well. Hence, by Lemma 14,

$$\|\mathbf{A} * \mathbf{B}\|_2^2 \leq \|(\mathbf{A}^T \mathbf{A}) * \mathbf{I}_m\|_2 \|\mathbf{I}_n * (\mathbf{B} \mathbf{B}^T)\|_2.$$

□

Therefore, if we write the partition of a matrix $\mathbf{Z}$ in terms of the non-overlapping partition given by the restriction and prolongation matrices $\mathbf{R}^{(i,-1)}$ and $\mathbf{P}^{(i,-1)}$ as

$$\mathbf{Z} = \sum_{i,j} \mathbf{P}^{(i,-1)} \left(\mathbf{R}^{(i,-1)} \mathbf{Z} \mathbf{P}^{(j,-1)} \right) \mathbf{R}^{(j,-1)},$$

we obtain

$$\begin{aligned} \rho(\mathbf{Z}^{*2}) &\leq \max_i \left\| \mathbf{R}^{(i,-1)} \mathbf{Z}^T \mathbf{Z} \mathbf{P}^{(i,-1)} \right\|_2^{1/2} \max_i \left\| \mathbf{R}^{(i,-1)} \mathbf{Z} \mathbf{Z}^T \mathbf{P}^{(i,-1)} \right\|_2^{1/2} \\ &= \max_i \left\| \mathbf{Z} \mathbf{P}^{(i,-1)} \right\|_2 \max_i \left\| \mathbf{Z}^T \mathbf{P}^{(i,-1)} \right\|_2. \end{aligned} \quad (8.3)$$

Lemma 32.

For all $u \in V_h$,

$$a(\mathcal{R}^{(i,-1)} u, \mathcal{R}^{(i,-1)} u) \leq Ch^{-1} a(u, u)$$

for some constant C that is independent of h and N .

The proof of Lemma 32 follows immediately from (8.1) in Lemma 24.

Proof of Lemma 26. Set $\mathbf{D} := \sum_i \mathbf{P}^{(i,-1)} \mathbf{A}^{(i,-1)} \mathbf{R}^{(i,-1)}$. We then find that

$$\begin{aligned} \|\mathbb{V}[\mathcal{E}_{CG}]\|_A &= \|\mathbf{N}_{CG}^{\otimes 2} \mathbb{V}[\mathcal{X}] \mathbf{A}^{\otimes 2}\|_A \\ &= \left\| \left(\mathbf{A}^{1/2} \mathbf{N}_{CG} \right)^{\otimes 2} \mathbb{V}[\mathcal{X}] \left(\mathbf{A}^{1/2} \right)^{\otimes 2} \right\|_2 \\ &= \left\| \left(\mathbf{A}^{1/2} \mathbf{N}_{CG} \mathbf{D}^{1/2} \right)^{\otimes 2} \mathbb{V}[\mathcal{X}] \left(\mathbf{D}^{-1/2} \mathbf{A}^{1/2} \right)^{\otimes 2} \right\|_2 \\ &= \rho \left(\mathbb{V}[\mathcal{X}]^{1/2} \left(\mathbf{D}^{1/2} \mathbf{N}_{CG} \mathbf{A} \mathbf{N}_{CG} \mathbf{D}^{1/2} \right)^{\otimes 2} \mathbb{V}[\mathcal{X}] \right. \\ &\quad \left. \left(\mathbf{D}^{-1/2} \mathbf{A} \mathbf{D}^{-1/2} \right)^{\otimes 2} \mathbb{V}[\mathcal{X}]^{1/2} \right)^{1/2} \end{aligned}$$

$$\begin{aligned}
&\leq \rho \left(\mathbb{V}[\mathcal{X}]^{1/2} \left(\mathbf{D}^{1/2} \mathbf{N}_{CG} \mathbf{D}^{1/2} \right)^{\otimes 2} \mathbb{V}[\mathcal{X}]^{1/2} \right)^{1/2} \\
&\quad \rho \left(\mathbb{V}[\mathcal{X}]^{1/2} \left(\mathbf{D}^{-1/2} \mathbf{A} \mathbf{D}^{-1/2} \right)^{\otimes 2} \mathbb{V}[\mathcal{X}]^{1/2} \right)^{1/2} \\
&= \varepsilon \rho \left(\left(\mathbf{D}^{1/2} \mathbf{N}_{CG} \mathbf{D}^{1/2} \right)^{*2} \right)^{1/2} \rho \left(\left(\mathbf{D}^{-1/2} \mathbf{A} \mathbf{D}^{-1/2} \right)^{*2} \right)^{1/2}.
\end{aligned}$$

Using the result of eq. (8.3) and the discrete partition of unity property

$$\mathbf{I} = \sum_j \mathbf{P}^{(j,-1)} \mathbf{R}^{(j,-1)},$$

one obtains that

$$\begin{aligned}
\rho \left(\left(\mathbf{D}^{-1/2} \mathbf{A} \mathbf{D}^{-1/2} \right)^{*2} \right) &\leq \max_i \left\| \mathbf{D}^{-1/2} \mathbf{A} \mathbf{D}^{-1/2} \mathbf{P}^{(i,-1)} \right\|_2^2 \\
&\leq \max_i \sum_j \left\| \mathbf{R}^{(j,-1)} \mathbf{D}^{-1/2} \mathbf{A} \mathbf{D}^{-1/2} \mathbf{P}^{(i,-1)} \right\|_2^2 \\
&= \max_i \sum_j \left\| \left(\mathbf{A}^{(j,-1)} \right)^{-1/2} \mathbf{R}^{(j,-1)} \mathbf{A} \mathbf{P}^{(i,-1)} \left(\mathbf{A}^{(i,-1)} \right)^{-1/2} \right\|_2^2 \\
&\leq \max_i \sum_j \epsilon_{ij}^2 \\
&\leq N_c.
\end{aligned}$$

By the same token,

$$\rho \left(\left(\mathbf{D}^{1/2} \mathbf{N}_{CG} \mathbf{D}^{1/2} \right)^{*2} \right) \leq \max_i \left\| \mathbf{D}^{1/2} \mathbf{N}_{CG} \mathbf{P}^{(i,-1)} \left(\mathbf{A}^{(i,-1)} \right)^{1/2} \right\|_2^2.$$

Since

$$\begin{aligned}
&\left\| \mathbf{D}^{1/2} \mathbf{N}_{CG} \mathbf{P}^{(i,-1)} \left(\mathbf{A}^{(i,-1)} \right)^{1/2} \right\|_2 \\
&\leq \left\| \mathbf{D}^{1/2} \mathbf{A}^{-1} \mathbf{P}^{(i,-1)} \left(\mathbf{A}^{(i,-1)} \right)^{1/2} \right\|_2 + \left\| \mathbf{D}^{1/2} (\mathbf{A}^{-1} - \mathbf{N}_{CG}) \mathbf{P}^{(i,-1)} \left(\mathbf{A}^{(i,-1)} \right)^{1/2} \right\|_2 \\
&\leq \left\| \mathbf{D}^{1/2} \mathbf{A}^{-1} \mathbf{P}^{(i,-1)} \left(\mathbf{A}^{(i,-1)} \right)^{1/2} \right\|_2 + \|\mathbf{D}\|_2 \|\mathbf{A}^{-1} - \mathbf{N}_{CG}\|_2 \\
&\leq \left\| \mathbf{D}^{1/2} \mathbf{A}^{-1} \mathbf{P}^{(i,-1)} \left(\mathbf{A}^{(i,-1)} \right)^{1/2} \right\|_2 + C,
\end{aligned}$$

we have

$$\rho \left(\left(\mathbf{D}^{1/2} \mathbf{N}_{CG} \mathbf{D}^{1/2} \right)^{*2} \right) \leq C + \max_i \sum_j \left\| \left(\mathbf{A}^{(j,-1)} \right)^{1/2} \mathbf{R}^{(j,-1)} \mathbf{A}^{-1} \mathbf{P}^{(i,-1)} \left(\mathbf{A}^{(i,-1)} \right)^{1/2} \right\|_2^2.$$

We expand the summand into

$$\begin{aligned} & \left\| \left(\mathbf{A}^{(j,-1)} \right)^{1/2} \mathbf{R}^{(j,-1)} \mathbf{A}^{-1} \mathbf{P}^{(i,-1)} \left(\mathbf{A}^{(i,-1)} \right)^{1/2} \right\|_2 \\ & \leq \left\| \left(\mathbf{A}^{(j,-1)} \right)^{1/2} \left(\mathbf{R}^{(j,-1)} \mathbf{A}^{-1} \mathbf{P}^{(j,-1)} \right)^{-1/2} \left(\mathbf{R}^{(j,-1)} \mathbf{A}^{-1} \mathbf{P}^{(j,-1)} \right)^{1/2} \mathbf{R}^{(j,-1)} \mathbf{A}^{-1} \mathbf{P}^{(i,-1)} \right. \\ & \quad \left. \left(\mathbf{R}^{(i,-1)} \mathbf{A}^{-1} \mathbf{P}^{(i,-1)} \right)^{1/2} \left(\mathbf{R}^{(i,-1)} \mathbf{A}^{-1} \mathbf{P}^{(i,-1)} \right)^{-1/2} \left(\mathbf{A}^{(i,-1)} \right)^{1/2} \right\|_2 \\ & \leq \rho \left(\mathbf{A}^{(j,-1)} \left(\mathbf{R}^{(j,-1)} \mathbf{A}^{-1} \mathbf{P}^{(j,-1)} \right)^{-1} \right)^{1/2} \rho \left(\mathbf{A}^{(i,-1)} \left(\mathbf{R}^{(i,-1)} \mathbf{A}^{-1} \mathbf{P}^{(i,-1)} \right)^{-1} \right)^{1/2} \\ & \quad \left\| \left(\mathbf{R}^{(j,-1)} \mathbf{A}^{-1} \mathbf{P}^{(j,-1)} \right)^{1/2} \mathbf{R}^{(j,-1)} \mathbf{A}^{-1} \mathbf{P}^{(i,-1)} \left(\mathbf{R}^{(i,-1)} \mathbf{A}^{-1} \mathbf{P}^{(i,-1)} \right)^{1/2} \right\|_2 \end{aligned}$$

Using that

$$\rho \left(\mathbf{A}^{(i,-1)} \left(\mathbf{R}^{(i,-1)} \mathbf{A}^{-1} \mathbf{P}^{(i,-1)} \right)^{-1} \right) = \sup_{u \in V_h} \frac{a(\mathcal{R}^{(i,-1)}u, \mathcal{R}^{(i,-1)}u)}{a(u, u)}$$

and that

$$\left\| \left(\mathbf{R}^{(j,-1)} \mathbf{A}^{-1} \mathbf{P}^{(j,-1)} \right)^{1/2} \mathbf{R}^{(j,-1)} \mathbf{A}^{-1} \mathbf{P}^{(i,-1)} \left(\mathbf{R}^{(i,-1)} \mathbf{A}^{-1} \mathbf{P}^{(i,-1)} \right)^{1/2} \right\|_2 \leq \eta_{ij}$$

we therefore obtain that

$$\rho \left(\left(\mathbf{D}^{1/2} \mathbf{N}_{CG} \mathbf{D}^{1/2} \right)^{*2} \right) \leq C + \left\{ \max_i \sup_{u \in V_h} \frac{a(\mathcal{R}^{(i,-1)}u, \mathcal{R}^{(i,-1)}u)}{a(u, u)} \right\}^2 \max_i \sum_j \eta_{ij}^2.$$

We conclude using Lemma 32. $\square$

While the last inequality in (8.2) is not sharp in general, we verify numerically that it describes the correct behavior for one dimensional problems. We calculate $\|\mathbb{V}[\mathcal{E}_{CG}]\|_A$ directly for varying mesh spacing h and number of partitions N . The resulting scaling curves are shown in Figure 8.5. This shows that the bound in Lemma 26 is sharp in 1D, and that we cannot conclude the convergence of the two level methods using the energy norm.

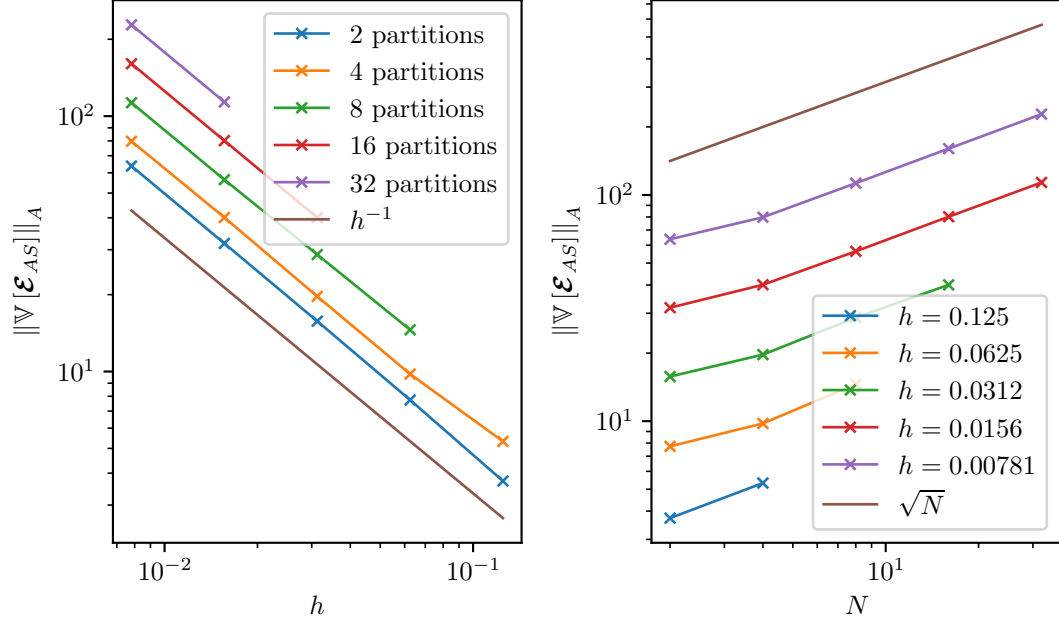


Figure 8.5: Energy norm of the variance of the fault-prone coarse-grid correction. Scaling with respect to mesh size h (left), and number of subdomains N (right).

On the other hand, if we consider the coarse grid correction on its own, we find that its Lyapunov spectral radius does not display *any* dependence on h or N :

Lemma 33.

If the prolongation in the coarse grid correction is protected, then it holds for the double energy norm $\|\bullet\|_{A^2} := \|\mathbf{A} \bullet \mathbf{A}^{-1}\|_2$, $\|\bullet\|_{A^2} := \|\mathbf{A}^{\otimes 2} \bullet (\mathbf{A}^{-1})^{\otimes 2}\|_2$ that

$$\|\mathbb{V}[\mathcal{E}_{CG}]\|_{A^2} \leq C\varepsilon$$

with C independent of ε , h and N , so that

$$\varrho(\mathcal{E}_{CG}) \leq \|\mathbf{E}_{CG}\|_2 + C\varepsilon.$$

Proof.

$$\begin{aligned} \|\mathbb{V}[\mathcal{E}_{CG}]\|_{A^2} &\leq \|\mathbf{N}_{CG}^{\otimes 2} \mathbb{V}[\mathcal{X}] \mathbf{A}^{\otimes 2}\|_{A^2} \\ &\leq \|\mathbf{A} \mathbf{N}_{CG}\|_2^2 \|\mathbb{V}[\mathcal{X}]\|_2 \end{aligned}$$

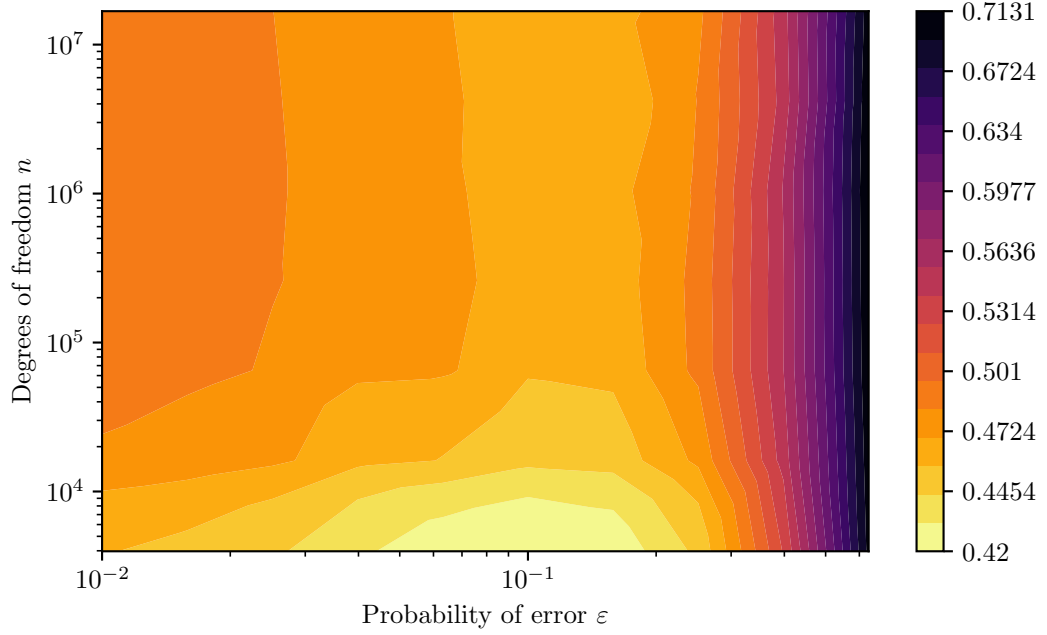


Figure 8.6: Asymptotic convergence rate $\varrho(\mathcal{E}_{ASH}\mathcal{E}_{CG})$ with protected prolongation and $N = 16$ subdomains.

$$\leq \varepsilon \|\mathbf{N}_{CG}\mathbf{A}\|_2^2.$$

Now, since $\|\mathbf{E}_{CG}\|_2 \leq C$ (see for example [62]), we have that $\|\mathbb{V}[\mathcal{E}_{CG}]\|_{A^2} \leq C\varepsilon$ and also that

$$\begin{aligned} \varrho(\mathcal{E}_{CG}) &\leq \sqrt{\|\mathbb{E}[\mathcal{E}_{CG}]\|_{A^2}^2 + \|\mathbb{V}[\mathcal{E}_{CG}]\|_{A^2}^2} \\ &\leq \|\mathbf{E}_{CG}\|_2 + C\varepsilon. \end{aligned}$$

□

While the coarse grid correction on itself is not a convergent iteration, this result shows that we can expect its fault-prone version to perform similarly to its fault-free variant. This is indeed what can be observed through numerical simulations, as seen in Figures 8.6 and 8.7. (The convergence is the fastest for $\varepsilon \approx 0.1$, not $\varepsilon = 0$. This indicates that some amount of damping might be beneficial for ASH.)

In order to obtain an analytic result for the full two level method, it would seem logical to analyze the domain decomposition solver in turn in the double energy

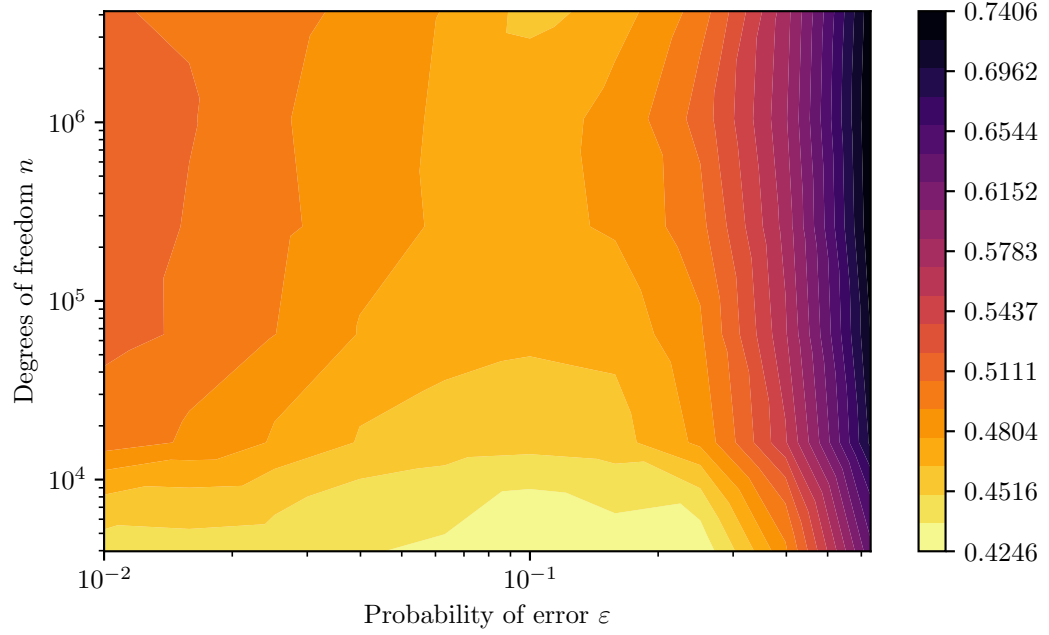


Figure 8.7: Asymptotic convergence rate $\rho(\mathcal{E}_{ASH}\mathcal{E}_{CG})$ with protected prolongation and $N \sim n$ subdomains.

or L^2 norm. The results in [25] suggest however that $\|\mathbf{N}_{AS}\mathbf{A}\|_{A^2} = \|\mathbf{N}_{AS}\mathbf{A}\|_2$ diverges as h goes to zero. It can be easily verified numerically (see Figure 8.8) that for a 1D test problem $\|\mathbf{N}_{AS}\mathbf{A}\|_2 \sim h^{-1/2}$, while $\rho(\mathbf{N}_{AS}\mathbf{A}) = \|\mathbf{N}_{AS}\mathbf{A}\|_A \leq C$.

8.5 Discussion

In this chapter, we generalized the framework to the analysis of fault-prone iterative schemes with overlap. We modeled blockwise faults, where special care had to be taken to account for the fact that affect unknowns can be owned by several subdomains at once. We gave convergence results for three commonly used one-level methods and observed in numerical examples that one-level solvers are resilient to faults, mirroring the behavior observed for the smoothing iterations of multigrid.

Inclusion of a coarse grid, as is mandatory for many applications, only results in a fault-resilient method, if the prolongation operation is protected. While the

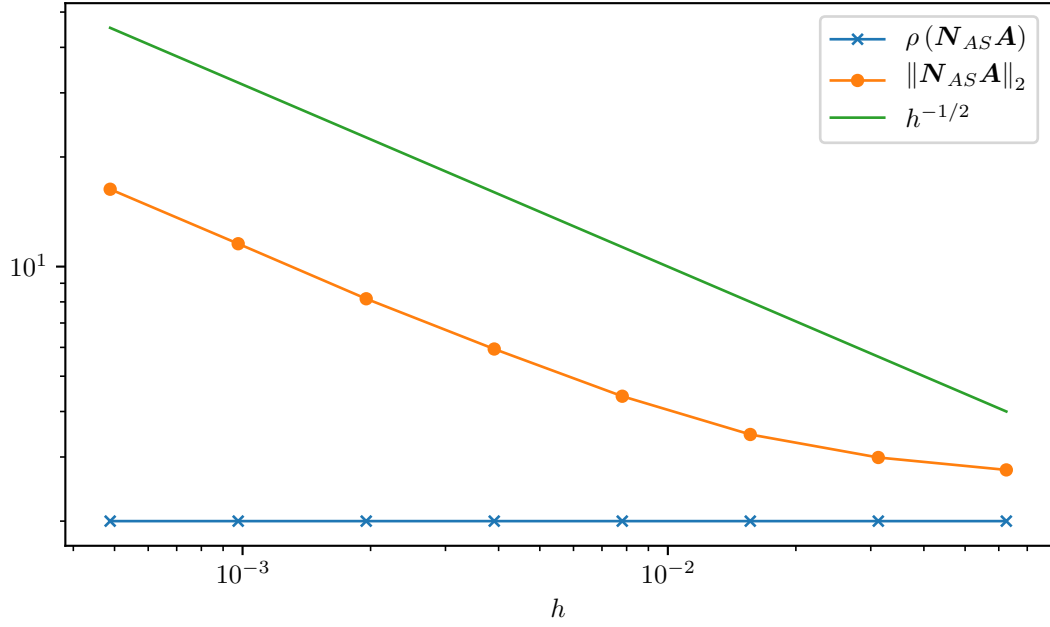


Figure 8.8: $\|N_{AS}A\|_2$ in 1D for several mesh sizes h and $N = 4$ subdomains.

numerical examples closely resemble the ones obtained for the multigrid method, the analysis of two-level methods is unsatisfactory. The issue that arises is that convergence analysis of the domain decomposition iteration can not be performed in spectral norm, contrary to multigrid. The isolated fault-prone coarse grid correction is non-divergent, as is the isolated domain decomposition iteration, but the analysis of their combination is inconclusive.

Chapter 9

Future Work on Fault Resilience and Related Topics

9.1 Fault-Resilient Domain Decomposition Methods

While the numerical experiments in Section 8.4 suggest that the same techniques as in the multigrid case can be used to obtain a fault-resilient two-level domain decomposition method, the application of the theoretical framework to this case is more intricate. The random faults matrices are best analyzed in L^2 -type norms, whereas domain decomposition methods are analyzed in energy norm. I believe that by a more tailored approach it is nevertheless possible to obtain convergence estimates for this important class of solvers. Equally well, it would be of interest to assert that domain decomposition methods without overlap (also called substructuring methods) can be made fault-resilient in the same way.

In Chapter 7, we discussed a simple approach to realize fault detection in all operations and protection of the prolongation. While the replication approach was chosen to illustrate the practical consequences of the theoretical results of the previous chapters, its implementation suffers from a large overhead. It might be prudent

to explore the use of node-level checksums for fault detection [68]. The computation of checksums would only introduce a negligible overhead, but come at the price of a coarser fault detection. Instead of applying the laissez-faire approach to a single faulty component, a whole block of values would be zeroed out. Moreover, protection of the prolongation might be achieved by using values from coarser levels of the hierarchy as implicit checkpoints.

Finally, there is interest to combine our work for linear iterative methods with the results obtained by other researchers for fault-prone Krylov methods and study the complete stack of linear systems solvers in high-performance computing.

9.2 Asynchronous Iterations

More generally, random matrices can be used not to model faults, but the delayed arrival of updates in an asynchronous iterative method. This touches on the challenge of communication costs in evermore parallel computation. Classically, iterative methods alternate phases of high communication volume between compute nodes with phases of high computational intensity. As supercomputers are built from more and more compute units, communication becomes a bottleneck for performance. Therefore, algorithms that are as asynchronous as possible will be required for the next generation of supercomputers. The study of such asynchronous iterative methods by itself is not a new topic [16, 54, 76, 91], but generally seems to be lacking sharp convergence estimates. In order to employ our novel approach based on random matrices, we would need to extend the theoretical framework to allow for memory terms, as to model the staggered arrival of updates.

Part II

Fast Solvers for Non-Local Equations

Chapter 10

The Finite Element Method for the Fractional Laplacian

In recent years, there has been a burgeoning of interest in the use of non-local and fractional models. To some extent, this move reflects the fact that with present day computational resources coupled with state of the art numerical algorithms, attention is now shifting back to the fidelity of the underlying mathematical models as opposed to their approximation. Fractional equations have been used to describe phenomena in anomalous diffusion [15], material science [13, 85], image processing [57, 75], finance [93] and electromagnetic fluids [77] [100]. Fractional order equations arise naturally as the limit of discrete diffusion governed by stochastic processes [79].

Whilst the development of fractional derivatives dates back to essentially the same time as their integer counterparts, the computational methods available for their numerical resolution drastically lags behind the vast array of numerical techniques from which one can choose to treat integer order partial differential equations. The recent literature abounds with work on numerical methods for fractional partial differential equations in one spatial dimension and fractional order temporal derivatives. However, with most applications of interest being posed on domains

in two or more spatial dimensions, the solution of fractional equations posed on complex domains is a problem of considerable practical interest.

The archetypal elliptic partial differential equation is the Poisson problem involving the standard Laplacian. By analogy, one can consider a fractional Poisson problem involving the fractional Laplacian. The first problem one encounters is that of how to define a fractional Laplacian, particularly in the case where the domain is compact, and a number of alternatives have been suggested. The *integral fractional Laplacian* is obtained by restriction of the Fourier definition to functions that have prescribed value outside of the domain of interest, whereas the spectral fractional Laplacian is based on the spectral decomposition of the regular Laplace operator. In general, the two operators are different [84], and only coincide when the domain of interest is the full space.

The approximation of the integral fractional Laplacian using finite elements was considered by D'Elia and Gunzburger [39]. The important work of Acosta and Borthagaray [2] and Grubb [61] gave regularity results for the analytic solution of the fractional Poisson problem and obtained convergence rates for the finite element approximation supported by numerical examples computed using techniques described in [1].

The numerical treatment of fractional partial differential equations is rather different from the integer order case owing to the fact that the fractional derivative is a non-local operator. This creates a number of issues including the fact that the resulting stiffness matrix is dense and, moreover, the entries in the matrix are given in terms of singular integrals. In turn, these features create issues in the numerical computation of the entries and the need to store the entries of a dense matrix, not to mention the fact that a solution of the resulting matrix equation has to be computed. The seasoned reader will readily appreciate that many of these issues are shared by boundary integral equations arising from classical integer order differential operators [86, 87, 103]. This similarity is not altogether surprising given

that the boundary integral operators are pseudo-differential operators of fractional order.

A different, integer order operator based approach, was taken by Nochetto, Otárola and Salgado [31, 80] for the case of the spectral Laplacian. Caffarelli and Silvestre [22] showed that the operator can be realized as a Dirichlet-to-Neumann operator of an extended problem in the half space in $d + 1$ dimensions.

In this part of the work, we explore the connection with boundary integral operators to develop techniques that enable one to efficiently tackle the integral fractional Laplacian. In particular, we develop techniques for the treatment of the stiffness matrix including the computation of the entries, the efficient storage of the resulting dense matrix and the efficient solution of the resulting equations. Since solutions of the fractional Poisson problem are generally non-smooth in a neighborhood of the boundary, we introduce a posteriori error estimators that are used for adaptive mesh refinement. The main ideas consist of generalizing proven techniques for the treatment of boundary integral equations to general fractional orders. Importantly, the approximation does not make any strong assumptions on the shape of the underlying domain and does not rely on any special structure of the matrix that could be exploited by fast transforms. We demonstrate the flexibility and performance of this approach in a couple of one- and two-dimensional numerical examples. It will be observed that optimal rates of convergence are achieved in all numerical examples.

10.1 Fractional Laplacians

The fractional Laplacian in $\mathbb{R}^d$ of order s , for $0 < s < 1$ and $d \in \mathbb{N}$, of a function u can be defined by the Fourier transform $\mathcal{F}$ as

$$(-\Delta)^s u = \mathcal{F}^{-1} [|\xi|^{2s} \mathcal{F}u] .$$

Alternatively, this expression can be rewritten [96] in integral form as

$$(-\Delta)^s u(\vec{x}) = C(d, s) \text{ p. v. } \int_{\mathbb{R}^d} d\vec{y} \frac{u(\vec{x}) - u(\vec{y})}{|\vec{x} - \vec{y}|^{d+2s}}$$

where

$$C(d, s) = \frac{2^{2s} s \Gamma(s + \frac{d}{2})}{\pi^{d/2} \Gamma(1 - s)}$$

is a normalization constant and p. v. denotes the Cauchy principal value of the integral [78, Chapter 5]. In the case where $s = 1$ this operator coincides with the usual Laplacian. If $\Omega \subset \mathbb{R}^d$ is a bounded Lipschitz domain, we define the *integral fractional Laplacian* $(-\Delta)^s$ to be the restriction of the full-space operator to functions with compact support in Ω . This generalizes the homogeneous Dirichlet condition applied in the case $s = 1$ to the case $s \in (0, 1)$.

The *fractional Poisson problem* is then given by

$$\begin{aligned} (-\Delta)^s u &= f \text{ in } \Omega, \\ u &= 0 \text{ in } \Omega^c \end{aligned} \tag{10.1}$$

This generalizes the homogeneous integer order Poisson problem from the case $s = 1$ to the case $s \in (0, 1)$.

A related operator on Ω is the *regional fractional Laplacian*

$$(-\Delta)_R^s u(\vec{x}) = C(d, s) \text{ p. v. } \int_{\Omega} d\vec{y} \frac{u(\vec{x}) - u(\vec{y})}{|\vec{x} - \vec{y}|^{d+2s}},$$

which generalizes the Laplacian with homogeneous Neumann boundary condition to the fractional order case. It will be seen that the techniques developed below for the integral fractional Laplacian also apply for the regional fractional Laplacian.

In practice, we are usually interested in dynamic behavior rather than steady state problems. The archetype of a time-dependent problem is the *fractional heat equation*

$$\partial_t u + (-\Delta)^s u = f \text{ in } [0, T] \times \Omega, \tag{10.2}$$

$$u = u_0 \text{ in } \{0\} \times \Omega.$$

Equally well, we might consider the fractional heat equation with the regional fractional Laplacian instead of the integral one, if a Neumann boundary condition is more appropriate.

10.2 Weak formulation

Define the usual fractional Sobolev space $H^s(\mathbb{R}^d)$ via the Fourier transform. If Ω is a sub-domain as above, then we define the Sobolev space $H^s(\Omega)$ to be

$$H^s(\Omega) := \left\{ u \in L^2(\Omega) \mid \|u\|_{H^s(\Omega)} < \infty \right\},$$

equipped with the norm

$$\|u\|_{H^s(\Omega)}^2 = \|u\|_{L^2(\Omega)}^2 + \int_{\Omega} d\vec{x} \int_{\Omega} d\vec{y} \frac{(u(\vec{x}) - u(\vec{y}))^2}{|\vec{x} - \vec{y}|^{d+2s}}.$$

The space

$$\tilde{H}^s(\Omega) := \{u \in H^s(\mathbb{R}^d) \mid u = 0 \text{ in } \Omega^c\}$$

can be equipped with the energy norm

$$\|u\|_{\tilde{H}^s(\Omega)} := a(u, u)^{1/2} = \sqrt{\frac{C(d, s)}{2}} |u|_{H^s(\mathbb{R}^d)},$$

where the non-standard factor $\sqrt{C(d, s)/2}$ is included for convenience. For $s > 1/2$, $\tilde{H}^s(\Omega)$ coincides with the space $H_0^s(\Omega)$ which is the closure of $C_0^\infty(\Omega)$ with respect to the $H^s(\Omega)$ -norm. For $s < 1/2$, $\tilde{H}^s(\Omega)$ is identical to $H^s(\Omega)$. In the critical case $s = 1/2$, $\tilde{H}^s(\Omega) \subset H_0^s(\Omega)$, and the inclusion is strict. (See for example [78, Chapter 3].)

The usual approach to dealing with elliptic PDEs consists of obtaining a weak form of the operator by multiplying the equation by a test function and applying

integration by parts [50]. In contrast, for equations involving the fractional Laplacian $(-\Delta)^s u$, we again multiply by as test function $v \in \tilde{H}^s(\Omega)$ and integrate over $\mathbb{R}^d$, and then, instead of integration by parts, we use the identity

$$\int_{\mathbb{R}^d} d\vec{x} \int_{\mathbb{R}^d} d\vec{y} \frac{(u(\vec{x}) - u(\vec{y})) v(\vec{x})}{|\vec{x} - \vec{y}|^{d+2s}} = - \int_{\mathbb{R}^d} d\vec{x} \int_{\mathbb{R}^d} d\vec{y} \frac{(u(\vec{x}) - u(\vec{y})) v(\vec{y})}{|\vec{x} - \vec{y}|^{d+2s}}.$$

Following this approach, since both u and v vanish outside of Ω , we arrive at the bilinear form

$$\begin{aligned} a(u, v) &= \frac{C(d, s)}{2} \int_{\Omega} d\vec{x} \int_{\Omega} d\vec{y} \frac{(u(\vec{x}) - u(\vec{y})) (v(\vec{x}) - v(\vec{y}))}{|\vec{x} - \vec{y}|^{d+2s}} \\ &\quad + C(d, s) \int_{\Omega} d\vec{x} \int_{\Omega^c} d\vec{y} \frac{u(\vec{x}) v(\vec{x})}{|\vec{x} - \vec{y}|^{d+2s}}, \end{aligned}$$

corresponding to $(-\Delta)^s$ on $\tilde{H}^s(\Omega) \times \tilde{H}^s(\Omega)$. The bilinear form $a(\cdot, \cdot)$ is trivially seen to be $\tilde{H}^s(\Omega)$ -coercive and continuous and, as such, is amenable to treatment using the Lax-Milgram Lemma.

In this chapter we shall concern ourselves with the computational details needed to implement the finite element approximation of problems involving the fractional Laplacian. To this end, the presence of the unbounded domain Ω^c in the bilinear form $a(\cdot, \cdot)$ is somewhat undesirable. Fortunately, we can dispense with Ω^c using the following argument. The identity

$$\frac{1}{|\vec{x} - \vec{y}|^{d+2s}} = \frac{1}{2s} \nabla_{\vec{y}} \cdot \frac{\vec{x} - \vec{y}}{|\vec{x} - \vec{y}|^{d+2s}},$$

enables the second integral to be rewritten using the Gauss theorem as

$$\frac{C(d, s)}{2s} \int_{\Omega} d\vec{x} \int_{\partial\Omega} d\vec{y} \frac{u(\vec{x}) v(\vec{x}) \vec{n}_{\vec{y}} \cdot (\vec{x} - \vec{y})}{|\vec{x} - \vec{y}|^{d+2s}},$$

where $\vec{n}_{\vec{y}}$ is the *inward* normal to $\partial\Omega$ at $\vec{y}$, so that the bilinear form can be expressed equivalently as

$$\begin{aligned} a(u, v) &= \frac{C(d, s)}{2} \int_{\Omega} d\vec{x} \int_{\Omega} d\vec{y} \frac{(u(\vec{x}) - u(\vec{y})) (v(\vec{x}) - v(\vec{y}))}{|\vec{x} - \vec{y}|^{d+2s}} \\ &\quad + \frac{C(d, s)}{2s} \int_{\Omega} d\vec{x} \int_{\partial\Omega} d\vec{y} \frac{u(\vec{x}) v(\vec{x}) \vec{n}_{\vec{y}} \cdot (\vec{x} - \vec{y})}{|\vec{x} - \vec{y}|^{d+2s}}. \end{aligned}$$

As an aside, we note that dropping the boundary term in $a(\cdot, \cdot)$ leads to the bilinear form

$$b(u, v) = \frac{C(d, s)}{2} \int_{\Omega} d\vec{x} \int_{\Omega} d\vec{y} \frac{(u(\vec{x}) - u(\vec{y}))(v(\vec{x}) - v(\vec{y}))}{|\vec{x} - \vec{y}|^{d+2s}}$$

which represents the regional fractional Laplacian [17, 32]. The regional fractional Laplacian can be interpreted as a generalization of the usual Laplacian with homogeneous Neumann boundary condition for $s = 1$ to the case of fractional orders $s \in (0, 1)$. It will transpire from our work that most of the presented techniques carry over to the regional fractional Laplacian by simply omitting the boundary integral terms.

10.3 Finite Element Approximation of the Fractional Poisson Equation

The fractional Poisson problem

$$\begin{aligned} (-\Delta)^s u &= f && \text{in } \Omega, \\ u &= 0 && \text{in } \Omega^c \end{aligned}$$

takes the variational form

$$\text{Find } u \in \tilde{H}^s(\Omega) : \quad a(u, v) = \langle f, v \rangle \quad \forall v \in \tilde{H}^s(\Omega). \quad (10.3)$$

The existence of a unique solution to the fractional Poisson problem eq. (10.1) (and its subsequent finite element approximation) follows from the Lax-Milgram Lemma.

As a special case of a result by Grubb [61], the regularity of the solution was recorded in [1]:

Theorem 34.

Let $\partial\Omega \in C^\infty$, $f \in H^r(\Omega)$ for $r \geq -s$ and $u \in \tilde{H}^s(\Omega)$ be the solution of the fractional

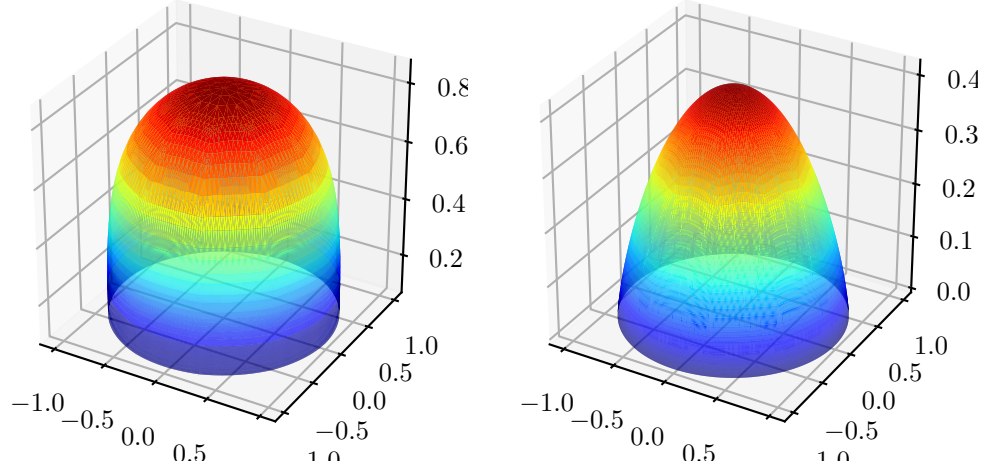


Figure 10.1: Solutions u corresponding to the constant right-hand side for $s = 0.25$ and for $s = 0.75$.

Poisson problem (10.3). Then the following regularity estimate holds:

$$u \in \begin{cases} H^{2s+r}(\Omega) & \text{if } 0 < s + r < 1/2, \\ H^{s+1/2-\varepsilon}(\Omega) \quad \forall \varepsilon > 0 & \text{if } 1/2 \leq s + r. \end{cases}$$

The result shows that increasing the regularity of the right-hand side only results in an increase of the regularity of the solution, as long as $r < 1/2 - s$. We contrast the result with the standard “lifting property” of the integer order Laplacian [50, Theorem 3.10], whereby a gain of two derivatives takes place on smooth domains. Examining a particular solution reveals why the lifting property does not extend to the fractional case. On $\Omega = B(0, 1) \subset \mathbb{R}^2$, the constant right-hand side $f = 2^{2s}\Gamma(1+s)^2$ has corresponding solution $u(\vec{x}) = (1 - |\vec{x}|^2)_+^s$, where $y_+ = \max(y, 0)$, and u is singular at the boundary [45, 58]. (See Figure 10.1.) It generally holds that the solutions of the fractional Poisson problem contain a power $\delta(\vec{x})^s \in H^{s+1/2-\varepsilon}$, where $\delta(\vec{x})$ is the distance of $\vec{x}$ to the boundary $\partial\Omega$.

For polygonal domains, the infinite lifting property does not hold in the integer order case either [50, Theorem 3.11].

Henceforth, let Ω be a polygon, and let $\mathcal{P}_h$ be a family of shape-regular and

locally quasi-uniform triangulations of Ω [50], and

$$\partial\mathcal{P}_h = \{\text{edge } e \mid \exists K \in \mathcal{P}_h : e \subset \partial K \cap \partial\Omega\}$$

the trace of the interior mesh. Let $\mathcal{N}_h$ be the set of vertices of $\mathcal{P}_h$ and h_K be the diameter of the element $K \in \mathcal{P}_h$, and h_e the diameter of $e \in \partial\mathcal{P}_h$. Moreover, let

$$\begin{aligned} h &:= \max_{K \in \mathcal{P}_h} h_K, \\ h_{\min} &:= \min_{K \in \mathcal{P}_h} h_K, \\ h_{\partial} &:= \max_{e \in \partial\mathcal{P}_h} h_e. \end{aligned}$$

Let ϕ_i be the usual piecewise linear basis function associated with a node $\vec{z}_i \in \mathcal{N}_h$, satisfying $\phi_i(\vec{z}_j) = \delta_{ij}$ for $\vec{z}_j \in \mathcal{N}_h$, and let $X_h := \text{span}\{\phi_i \mid \vec{z}_i \in \mathcal{N}_h\}$. The finite element subspace $V_h \subset \tilde{H}^s(\Omega)$ is given by $V_h = X_h$ when $s < 1/2$ and by

$$V_h = \{v_h \in X_h \mid v_h = 0 \text{ on } \partial\Omega\} = \text{span}\{\phi_i \mid \vec{z}_i \notin \partial\Omega\}$$

when $s \geq 1/2$. The corresponding set of degrees of freedom $\mathcal{I}_h$ for V_h is given by $\mathcal{I}_h = \mathcal{N}_h$ when $s < 1/2$ and otherwise consists of nodes in the interior of Ω . In both cases we denote the cardinality of $\mathcal{I}_h$ by n . The set of degrees of freedom on an element $K \in \mathcal{P}_h$ is denoted by $\mathcal{I}_K$.

The stiffness matrix associated with the fractional Laplacian is defined to be $\mathbf{A}^s = \{a(\phi_i, \phi_j)\}_{i,j}$, where

$$\begin{aligned} a(\phi_i, \phi_j) &= \frac{C(d, s)}{2} \int_{\Omega} d\vec{x} \int_{\Omega} d\vec{y} \frac{(\phi_i(\vec{x}) - \phi_i(\vec{y}))(\phi_j(\vec{x}) - \phi_j(\vec{y}))}{|\vec{x} - \vec{y}|^{d+2s}} \\ &\quad + \frac{C(d, s)}{2s} \int_{\Omega} d\vec{x} \int_{\partial\Omega} d\vec{y} \frac{\phi_i(\vec{x}) \phi_j(\vec{x}) \vec{n}_{\vec{y}} \cdot (\vec{x} - \vec{y})}{|\vec{x} - \vec{y}|^{d+2s}}. \end{aligned}$$

Let $u_h \in V_h$ be the finite element solution. By Céa's Lemma and the inclusion $\tilde{H}^s(\Omega) \subseteq H^s(\Omega)$,

$$\|u - u_h\|_{\tilde{H}^s(\Omega)} \leq C \inf_{v_h \in V_h} \|u - v_h\|_{\tilde{H}^s(\Omega)} \leq C \inf_{v_h \in V_h} \|u - v_h\|_{H^s(\Omega)}. \quad (10.4)$$

The proof of the rate of convergence is based on approximation properties of the Scott-Zhang interpolation operator with respect to the $\|\cdot\|_{H^s(\Omega)}$ -norm. For $u \in H^\ell(\Omega)$, $\ell > 1/2$, let

$$\Pi_h u = \sum_i \left(\int_{K_i} \psi_i(\vec{y}) u(\vec{y}) d\vec{y} \right) \phi_i$$

denote the Scott-Zhang interpolation operator [83]. Here, K_i is either a simplex K of the triangulation such that $i \in \mathcal{I}_K$, or a sub-simplex e such that i is associated with e , and $\{\psi_i\}_i$ is an $L^2(K_i)$ -dual basis chosen such that $\int_{K_i} \psi_i(\vec{y}) \phi_j(\vec{y}) d\vec{y} = \delta_{ij}$.

Acosta and Borthagaray [2] showed that if $u \in H^\ell(S_K)$, where $S_K = \bigcup_{\tilde{K} \text{ s.t. } \bar{K} \cap \bar{\tilde{K}} \neq \emptyset} \tilde{K}$ is the patch of an element $K \in \mathcal{P}_h$, the local approximation property

$$\int_K d\vec{x} \int_{S_K} d\vec{y} \frac{|(u - \Pi_h u)(\vec{x}) - (u - \Pi_h u)(\vec{y})|^2}{|\vec{x} - \vec{y}|^{d+2s}} \leq C h_K^{2\ell-2s} |u|_{H^\ell(S_K)}^2$$

holds for $0 < s < \ell < 1$ and for $1/2 < s < 1$ and $1 < \ell < 2$. Below, we extend the result to include the case of $0 < s < 1/2$ and $1 \leq \ell \leq 2$ as well. This result will be needed to take advantage of higher regularity of the solution in the interior of the domain.

Lemma 35.

If $u \in H^\ell(S_K)$, for $\ell \in (1/2, 2]$ and $0 < s \leq \ell$, then

$$\int_K d\vec{x} \int_{S_K} d\vec{y} \frac{|\Pi_h u(\vec{x}) - \Pi_h u(\vec{y})|^2}{|\vec{x} - \vec{y}|^{d+2s}} \leq C \left[h_K^{-2s} \|u\|_{L^2(S_K)}^2 + h_K^{2\ell-2s} |u|_{H^\ell(S_K)}^2 \right], \quad (10.5)$$

$$\int_K d\vec{x} \int_{S_K} d\vec{y} \frac{|(u - \Pi_h u)(\vec{x}) - (u - \Pi_h u)(\vec{y})|^2}{|\vec{x} - \vec{y}|^{d+2s}} \leq C h_K^{2\ell-2s} |u|_{H^\ell(S_K)}^2. \quad (10.6)$$

Proof. Similar to the proof in [36], for $t \in (1/2, \min\{1, \ell\}]$, we obtain

$$\int_K d\vec{x} \int_{S_K} d\vec{y} \frac{|\Pi_h u(\vec{x}) - \Pi_h u(\vec{y})|^2}{|\vec{x} - \vec{y}|^{d+2s}} \leq C \sum_{i \in \mathcal{I}_K} \|\psi_i\|_{L^\infty(K_i)}^2 \|u\|_{L^1(K_i)}^2 |\phi_i|_{H^s(S_K)}^2.$$

From [10, Theorem 4.8], we obtain

$$|\phi_i|_{H^s(S_K)}^2 \leq C h_K^{d-2s},$$

and from [83, Lemma 3.1] that

$$\|\psi_i\|_{L^\infty(K_i)}^2 \leq Ch_K^{-2 \dim K_i}.$$

If K_i is a simplex K , then

$$\|u\|_{L^1(K_i)}^2 \leq h_K^d \|u\|_{L^2(S_K)}^2.$$

On the other hand, if K_i is a sub-simplex e , then by [36, 83]

$$\|u\|_{L^1(K_i)}^2 \leq C \left(h_K^{d-2} \|u\|_{L^2(S_K)}^2 + h_K^{d-2+2t} |u|_{H^t(S_K)}^2 \right).$$

Combining these estimates gives the following stability estimate

$$\int_K d\vec{x} \int_{S_K} d\vec{y} \frac{|\Pi_h u(\vec{x}) - \Pi_h u(\vec{y})|^2}{|\vec{x} - \vec{y}|^{d+2s}} \leq C \left[h_K^{-2s} \|u\|_{L^2(S_K)}^2 + h_K^{2t-2s} |u|_{H^t(S_K)}^2 \right].$$

Now, let p be a polynomial of degree 1 on S_K . Then, by $\Pi_h p = p$ and the above stability result

$$\begin{aligned} & \int_K d\vec{x} \int_{S_K} d\vec{y} \frac{|(u - \Pi_h u)(\vec{x}) - (u - \Pi_h u)(\vec{y})|^2}{|\vec{x} - \vec{y}|^{d+2s}} \\ & \leq C \left[\int_K d\vec{x} \int_{S_K} d\vec{y} \frac{|(u - p)(\vec{x}) - (u - p)(\vec{y})|^2}{|\vec{x} - \vec{y}|^{d+2s}} \right. \\ & \quad \left. + \int_K d\vec{x} \int_{S_K} d\vec{y} \frac{|(\Pi_h p - \Pi_h u)(\vec{x}) - (\Pi_h p - \Pi_h u)(\vec{y})|^2}{|\vec{x} - \vec{y}|^{d+2s}} \right] \\ & \leq C \left[|u - p|_{H^s(S_K)}^2 + h_K^{-2s} \|u - p\|_{L^2(S_K)}^2 + h_K^{2t-2s} |u - p|_{H^t(S_K)}^2 \right]. \end{aligned}$$

We distinguish two cases. In the first case, assume that $\ell < 1$ and take $t = \ell$. By the Bramble-Hilbert Lemma, it holds that there exists a polynomial p of order zero such that $\|u - p\|_{L^2(S_K)} \leq Ch_K^\ell |u|_{H^\ell(S_K)}$. Moreover $|u - p|_{H^\ell(S_K)} = |u|_{H^\ell(S_K)}$ and by interpolation of spaces $|u - p|_{H^s(S_K)} \leq Ch_K^{\ell-s} |u|_{H^\ell(S_K)}$. Hence

$$\int_K d\vec{x} \int_{S_K} d\vec{y} \frac{|(u - \Pi_h u)(\vec{x}) - (u - \Pi_h u)(\vec{y})|^2}{|\vec{x} - \vec{y}|^{d+2s}} \leq Ch_K^{2\ell-2s} |u|_{H^\ell(S_K)}^2. \quad (10.7)$$

In the second case, assume that $\ell \in [1, 2]$, and take $t = 1$. Again, by the Bramble-Hilbert Lemma, a polynomial p of degree one can be found such that

$\|u - p\|_{L^2(S_K)} \leq Ch_K^\ell |u|_{H^\ell(S_K)}$ and $|u - p|_{H^1(S_K)} \leq Ch_K^{\ell-1} |u|_{H^1(S_K)}$. By interpolation of spaces, it is obtained that $|u - p|_{H^s(S_K)} \leq Ch_K^{\ell-s} |u|_{H^s(S_K)}$, and therefore the local approximability result (10.7) holds in this case as well. $\square$

The proof of the rate of convergence of the finite element solution is an immediate consequence of Lemma 35.

Lemma 36.

If $u \in H^t(\Omega) \cap H_{\text{loc}}^\ell(\Omega)$, for $t, \ell \in (1/2, 2]$ and $0 < s \leq t \leq \ell$, i.e. if u has Sobolev regularity t and interior regularity ℓ , then

$$\|u - u_h\|_{\tilde{H}^s(\Omega)} \leq C \left(h^{\ell-s} |u|_{H_{\text{loc}}^\ell(\Omega)} + h_\partial^{t-s} |u|_{H^t(\Omega)} \right), \quad (10.8)$$

where h_∂ is the maximum size of all elements K whose patch S_K touches the boundary. In particular, if the family of triangulations $\mathcal{P}_h$ is globally quasi-uniform, and $u \in H^t(\Omega)$, for $t \in (1/2, 2]$ and $0 < s \leq t$, then

$$\|u - u_h\|_{\tilde{H}^s(\Omega)} \leq Ch^{t-s} |u|_{H^t(\Omega)}. \quad (10.9)$$

Proof. Assume that $u \in H^t(\Omega) \cap H_{\text{loc}}^\ell(\Omega)$, for $t, \ell \in (1/2, 2]$ and $0 < s \leq t \leq \ell$. By a localization result of Faermann [51, 52], we can estimate

$$\begin{aligned} & \|u - \Pi_h u\|_{H^s(\Omega)}^2 \\ & \leq C \sum_K \left[\int_K d\vec{x} \int_{S_K} d\vec{y} \frac{|(u - \Pi_h u)(\vec{x}) - (u - \Pi_h u)(\vec{y})|^2}{|\vec{x} - \vec{y}|^{d+2s}} + h_K^{-2s} \|u - \Pi_h u\|_{L^2(K)}^2 \right]. \end{aligned} \quad (10.10)$$

The local approximability of the Scott-Zhang Interpolation (10.6) then gives

$$\begin{aligned} \|u - \Pi_h u\|_{H^s(\Omega)}^2 & \leq C \left\{ \sum_{\substack{K \\ S_K \cap \partial\Omega = \emptyset}} h_K^{2\ell-2s} |u|_{H^\ell(S_K)}^2 + \sum_{\substack{K \\ S_K \cap \partial\Omega \neq \emptyset}} h_K^{2t-2s} |u|_{H^t(S_K)}^2 \right\} \\ & \leq C \left(h^{2\ell-2s} |u|_{H_{\text{loc}}^\ell(\Omega)}^2 + h_\partial^{2t-2s} |u|_{H^t(\Omega)}^2 \right), \end{aligned}$$

so that we can conclude using eq. (10.4). $\square$

Moreover, using a standard Aubin-Nitsche argument [50, Lemma 2.31] gives estimates in $L^2(\Omega)$:

Theorem 37 ([18]).

If the family of triangulations $\mathcal{P}_h$ is shape regular and globally quasi-uniform, and, for $\epsilon > 0$, $u \in H^{s+1/2-\epsilon}(\Omega)$, then

$$\|u - u_h\|_{L^2} \leq \begin{cases} C(s, \epsilon) h^{1/2+s-\epsilon} |u|_{H^{s+1/2-\epsilon}(\Omega)} & \text{if } 0 < s < 1/2, \\ C(s, \epsilon) h^{1-2\epsilon} |u|_{H^{s+1/2-\epsilon}(\Omega)} & \text{if } 1/2 \leq s < 1. \end{cases}$$

Estimate (10.9) implies that if $u \in H^2(\Omega)$, then the expected rate of convergence on a globally quasi-uniform mesh is $h^{2-s} \sim n^{(s-2)/d}$. However, the solutions of (10.1) generally have limited regularity as described in Theorem 34, owing to the lack of regularity in the neighborhood of the boundary. This suggests one should use a more finely graded mesh near the boundary so that $h_{\partial} \ll h$. The estimate (10.8) distinguishes between elements in the interior of the domain Ω and elements touching the boundary $\partial\Omega$. Generally, one would hope to restore the optimal rate of convergence $\mathcal{O}(n^{(s-2)/d})$ observed for smooth solutions by using appropriately graded meshes.

To this end, we choose $h_{\partial} \sim h^{(\ell-s)/(1-2\epsilon)}$, where h denotes the maximum element size in the interior. Applying estimate (10.8) with $t = s + 1/2 - \epsilon$ we obtain

$$\|u - u_h\|_{\tilde{H}^s(\Omega)} \leq Ch^{\ell-s} \left(|u|_{H_{\text{loc}}^{\ell}(\Omega)} + |u|_{H^{1/2+s-\epsilon}(\Omega)} \right).$$

In the one-dimensional case, $d = 1$, it is a simple matter to construct meshes for which $h_K = h$ for all elements K such that $S_K \cap \partial\Omega = \emptyset$, and $h_K = h^{(\ell-s)/(1-2\epsilon)}$ for the four elements whose patches touch the boundary. Then the number of unknowns $n \sim h^{-1}$ and therefore

$$\|u - u_h\|_{\tilde{H}^s(\Omega)} \leq Cn^{s-\ell} \left(|u|_{H_{\text{loc}}^{\ell}(\Omega)} + |u|_{H^{1/2+s-\epsilon}(\Omega)} \right). \quad (10.11)$$

This means that, provided $u \in H_{\text{loc}}^2(\Omega)$, the optimal rate of convergence is indeed achieved.

However, in the two-dimensional case one is constrained in the construction of the meshes if, as we do here, the use of hanging nodes is prohibited. This raises the possibility of the optimal rate of convergence of $\mathcal{O}(n^{(s-2)/d})$ not being achievable in practice. Acosta and Borthagaray [2] proposed a graded mesh for the unit disc domain for approximating the solution obtained with a constant right-hand side, in which a grading parameter $\mu \geq 1$ is selected and elements are distributed in M concentric layers of radii $r_i = 1 - \left(1 - \frac{i}{M}\right)^\mu$, $i = 1, \dots, M$, so that the element sizes are $h_i = r_i - r_{i-1}$. The dimension of the associated finite element space is $n \sim M^2$ if $\mu \leq 2$, and $n \sim M^\mu$ if $\mu > 2$. By following the arguments in [2] it is shown that

$$\|u - u_h\|_{\tilde{H}^s(\Omega)} \leq C n^{(\max\{s-\ell, -1+\varepsilon\})/2} \left(|u|_{H_{\text{loc}}^\ell(\Omega)} + |u|_{H^{1/2+s-\varepsilon}(\Omega)} \right). \quad (10.12)$$

This suggests that the optimal rate in $d = 2$ dimensions is $n^{-1/2+\varepsilon}$, which can be achieved if u has interior regularity of at least order $1 + s - \varepsilon$. We note that contrary to the result stated in [2], this holds for all $s \in (0, 1)$, provided that the solution u enjoys sufficient interior regularity.

A priori mesh grading of the type discussed above can give considerable improvements but does rely on the availability of sharp estimates for the regularity which, unfortunately, is not always the case for practical problems on complicated domains. Therefore, we explore a posteriori error indicators for the fractional Laplacian in Section 10.4.

10.4 Adaptive Mesh Refinement

The fractional Laplacian $(-\Delta)^s : \tilde{H}^s(\Omega) \rightarrow H^{-s}(\Omega)$ is continuous with continuous inverse. Hence, the discretization error

$$e = u - u_h$$

can be bounded from above and below in terms of the residual $r = f - (-\Delta)^s u_h$ as

$$c \|r\|_{H^{-s}(\Omega)} \leq \|e\|_{\tilde{H}^s(\Omega)} \leq C \|r\|_{H^{-s}(\Omega)},$$

where c and C are positive constants depending on the norm of $(-\Delta)^s$ and its inverse. It was shown in [51, 52] that the Babuška-Rheinboldt type error indicators

$$\eta_i := \sup_{\substack{v \in \tilde{H}^s(\Omega) \\ v\phi_i \neq 0}} \frac{|a(e, v\phi_i)|}{\|v\phi_i\|_{\tilde{H}^s(\Omega)}}$$

are reliable and efficient for non-negative order differential operators. Unfortunately, the estimators are not computable owing to the presence of the supremum. However, η_i can be estimated up to a generic (unknown) factor, independent of the solution and the mesh-size, by

$$\hat{\eta}_i := \sqrt{\sum_{K \in S_i} h_K^{2s} \|r\|_{L^2(K)}^2}, \quad (10.13)$$

where S_i is the patch formed by the elements which contain the node $\vec{z}_i$. Therefore, the following a posteriori estimate holds:

$$\|e\|_{\tilde{H}^s(\Omega)}^2 \leq \hat{\eta} := \sum_i \hat{\eta}_i^2.$$

The estimator (10.13) involves only the local L^2 -norms of the residual and can easily be evaluated using quadrature, but then one requires pointwise values for $(-\Delta)^s u_h$. A direct evaluation using quadrature rules is discouraged due to the presence of the singularity of the kernel $|\vec{x} - \vec{y}|^{-d-2s}$ as well as the fact that the domain of integration is unbounded. The next result gives a method whereby this issue can be circumvented.

Lemma 38.

For $u \in V_h$ and $\vec{x} \in K_0$ for an element $K_0 \in \mathcal{P}_h$, the fractional Laplacian $(-\Delta)^s u$ can be evaluated as

$$\begin{aligned} \frac{(-\Delta)^s u(\vec{x})}{C(d, s)} &= \frac{1}{d+2s-2} \int_{\partial K_0} \frac{\nabla u|_{K_0} \cdot \vec{n}_{\vec{y}}}{|\vec{x} - \vec{y}|^{d+2s-2}} d\vec{y} - \frac{u(\vec{x})}{2s} \int_{\partial K_0} \frac{\vec{n}_{\vec{y}} \cdot (\vec{x} - \vec{y})}{|\vec{x} - \vec{y}|^{d+2s}} d\vec{y} \\ &\quad + \sum_{K \neq K_0} \left\{ \frac{1}{2s(d+2s-2)} \int_{\partial K} \nabla u|_K \cdot \vec{n}_{\vec{y}} \frac{1}{|\vec{x} - \vec{y}|^{d+2s-2}} d\vec{y} \right\} \end{aligned} \quad (10.14)$$

$$\left. -\frac{1}{2s} \int_{\partial K} u(\vec{y}) \frac{\vec{n}_{\vec{y}} \cdot (\vec{x} - \vec{y})}{|\vec{x} - \vec{y}|^{d+2s}} d\vec{y} \right\} \quad (10.15)$$

if $s \neq 1 - d/2$, and as

$$\begin{aligned} \frac{(-\Delta)^s u(\vec{x})}{C(d, s)} &= \int_{\partial K_0} \nabla u|_{K_0} \cdot \vec{n}_{\vec{y}} \log \frac{1}{|\vec{x} - \vec{y}|} d\vec{y} - \frac{u(\vec{x})}{2s} \int_{\partial K_0} \frac{\vec{n}_{\vec{y}} \cdot (\vec{x} - \vec{y})}{|\vec{x} - \vec{y}|^{d+2s}} d\vec{y} \\ &\quad + \sum_{K \neq K_0} \left\{ \frac{1}{2s} \int_{\partial K} \nabla u|_K \cdot \vec{n}_{\vec{y}} \log \frac{1}{|\vec{x} - \vec{y}|} d\vec{y} - \frac{1}{2s} \int_{\partial K} u(\vec{y}) \frac{\vec{n}_{\vec{y}} \cdot (\vec{x} - \vec{y})}{|\vec{x} - \vec{y}|^2} d\vec{y} \right\}. \end{aligned} \quad (10.16)$$

when $s = 1 - d/2$.

Proof. Let $u \in V_h$ and $\vec{x} \in K_0$, for some $K_0 \in \mathcal{P}_h$. Then

$$\begin{aligned} \frac{(-\Delta)^s u(\vec{x})}{C(d, s)} &= \text{p. v.} \int_{\mathbb{R}^d} \frac{u(\vec{x}) - u(\vec{y})}{|\vec{x} - \vec{y}|^{d+2s}} d\vec{y} \\ &= \text{p. v.} \int_{K_0} \frac{u(\vec{x}) - u(\vec{y})}{|\vec{x} - \vec{y}|^{d+2s}} d\vec{y} + u(\vec{x}) \int_{\mathbb{R}^d \setminus K_0} \frac{1}{|\vec{x} - \vec{y}|^{d+2s}} d\vec{y} \\ &\quad - \sum_{K \neq K_0} \int_K \frac{u(\vec{y})}{|\vec{x} - \vec{y}|^{d+2s}} d\vec{y}. \end{aligned} \quad (10.17)$$

We have the following helpful identities:

$$\frac{\vec{x} - \vec{y}}{|\vec{x} - \vec{y}|^{d+2s}} = \begin{cases} \frac{1}{d+2s-2} \nabla_{\vec{y}} \frac{1}{|\vec{x} - \vec{y}|^{d+2s-2}}, & s \neq 1 - d/2, \\ \nabla_{\vec{y}} \log \frac{1}{|\vec{x} - \vec{y}|}, & s = 1 - d/2, \end{cases} \quad (10.18)$$

$$\frac{1}{|\vec{x} - \vec{y}|^{d+2s}} = \frac{1}{2s} \nabla_{\vec{y}} \cdot \frac{\vec{x} - \vec{y}}{|\vec{x} - \vec{y}|^{d+2s}}, \quad (10.19)$$

$$\frac{1}{|\vec{x} - \vec{y}|^{d+2s}} = \begin{cases} \frac{1}{2s(d+2s-2)} \Delta_{\vec{y}} \frac{1}{|\vec{x} - \vec{y}|^{d+2s-2}}, & s \neq 1 - d/2, \\ \frac{1}{d-2} \Delta_{\vec{y}} \log \frac{1}{|\vec{x} - \vec{y}|}, & s = 1 - d/2. \end{cases} \quad (10.20)$$

For $\vec{x}, \vec{y} \in K_0$,

$$u(\vec{x}) - u(\vec{y}) = \nabla u|_{K_0} \cdot (\vec{x} - \vec{y})$$

since $u|_{K_0}$ is linear. Hence, the first term of (10.17) can be rewritten using (10.18)

as

$$\text{p. v.} \int_{K_0} \frac{u(\vec{x}) - u(\vec{y})}{|\vec{x} - \vec{y}|^{d+2s}} d\vec{y} = \lim_{\varepsilon \rightarrow 0} \nabla u|_{K_0} \cdot \int_{K_0 \setminus B(\vec{x}, \varepsilon)} \frac{\vec{x} - \vec{y}}{|\vec{x} - \vec{y}|^{d+2s}} d\vec{y}$$

$$\begin{aligned}
&= \lim_{\varepsilon \rightarrow 0} \left[\nabla u|_{K_0} \cdot \frac{1}{d+2s-2} \int_{\partial K_0} \frac{\vec{n}_{\vec{y}}}{|\vec{x} - \vec{y}|^{d+2s-2}} d\vec{y} \right. \\
&\quad \left. + \nabla u|_{K_0} \cdot \frac{1}{d+2s-2} \underbrace{\int_{\partial B(\vec{x}, \varepsilon)} \frac{\vec{n}_{\vec{y}}}{|\vec{x} - \vec{y}|^{d+2s-2}} d\vec{y}}_{=0} \right] \\
&= \frac{1}{d+2s-2} \int_{\partial K_0} \frac{\nabla u|_{K_0} \cdot \vec{n}_{\vec{y}}}{|\vec{x} - \vec{y}|^{d+2s-2}} d\vec{y}.
\end{aligned}$$

Here, $\vec{n}_{\vec{y}}$ is the outward normal at $\vec{y}$. When $s = 1 - d/2$, we obtain instead that

$$\text{p. v.} \int_{K_0} \frac{u(\vec{x}) - u(\vec{y})}{|\vec{x} - \vec{y}|^2} d\vec{y} = \int_{\partial K_0} \nabla u|_{K_0} \cdot \vec{n}_{\vec{y}} \log \frac{1}{|\vec{x} - \vec{y}|} d\vec{y}.$$

The second term of (10.17) can be rewritten using (10.19) as

$$\begin{aligned}
\int_{\mathbb{R}^d \setminus K_0} \frac{1}{|\vec{x} - \vec{y}|^{d+2s}} d\vec{y} &= \frac{1}{2s} \int_{\mathbb{R}^d \setminus K_0} \nabla_{\vec{y}} \frac{\vec{x} - \vec{y}}{|\vec{x} - \vec{y}|^{d+2s}} d\vec{y} \\
&= -\frac{1}{2s} \int_{\partial K_0} \frac{\vec{n}_{\vec{y}} \cdot (\vec{x} - \vec{y})}{|\vec{x} - \vec{y}|^{d+2s}} d\vec{y}.
\end{aligned}$$

Finally, for every $K \in \mathcal{P}_h$, $K \neq K_0$, it follows from (10.20), Green's second identity and (10.18) that

$$\begin{aligned}
\int_K \frac{u(\vec{y})}{|\vec{x} - \vec{y}|^{d+2s}} d\vec{y} &= \frac{1}{2s(d+2s-2)} \int_K u(\vec{y}) \Delta_{\vec{y}} \frac{1}{|\vec{x} - \vec{y}|^{d+2s-2}} d\vec{y} \\
&= \frac{1}{2s(d+2s-2)} \int_K \underbrace{(\Delta_{\vec{y}} u(\vec{y}))}_{=0} \frac{1}{|\vec{x} - \vec{y}|^{d+2s-2}} d\vec{y} \\
&\quad + \frac{1}{2s(d+2s-2)} \int_{\partial K} \left\{ u(\vec{y}) (d+2s-2) \frac{\vec{n}_{\vec{y}} \cdot (\vec{x} - \vec{y})}{|\vec{x} - \vec{y}|^{d+2s}} \right. \\
&\quad \left. - \nabla u|_K \cdot \vec{n}_{\vec{y}} \frac{1}{|\vec{x} - \vec{y}|^{d+2s-2}} \right\} d\vec{y} \\
&= \frac{1}{2s} \int_{\partial K} u(\vec{y}) \frac{\vec{n}_{\vec{y}} \cdot (\vec{x} - \vec{y})}{|\vec{x} - \vec{y}|^{d+2s}} d\vec{y} \\
&\quad - \frac{1}{2s(d+2s-2)} \int_{\partial K} \nabla u|_K \cdot \vec{n}_{\vec{y}} \frac{1}{|\vec{x} - \vec{y}|^{d+2s-2}} d\vec{y}.
\end{aligned}$$

When $s = 1 - d/2$, we obtain instead that

$$\int_K \frac{u(\vec{y})}{|\vec{x} - \vec{y}|^{d+2s}} d\vec{y} = \frac{1}{2s} \int_{\partial K} u(\vec{y}) \frac{\vec{n}_{\vec{y}} \cdot (\vec{x} - \vec{y})}{|\vec{x} - \vec{y}|^2} d\vec{y} - \frac{1}{2s} \int_{\partial K} \nabla u|_K \cdot \vec{n}_{\vec{y}} \log \frac{1}{|\vec{x} - \vec{y}|} d\vec{y}.$$

In conclusion, we find that

$$\begin{aligned} \frac{(-\Delta)^s u(\vec{x})}{C(d, s)} &= \frac{1}{d + 2s - 2} \int_{\partial K_0} \frac{\nabla u|_{K_0} \cdot \vec{n}_{\vec{y}}}{|\vec{x} - \vec{y}|^{d+2s-2}} d\vec{y} - \frac{u(\vec{x})}{2s} \int_{\partial K_0} \frac{\vec{n}_{\vec{y}} \cdot (\vec{x} - \vec{y})}{|\vec{x} - \vec{y}|^{d+2s}} d\vec{y} \\ &\quad + \sum_{K \neq K_0} \left\{ \frac{1}{2s(d + 2s - 2)} \int_{\partial K} \nabla u|_K \cdot \vec{n}_{\vec{y}} \frac{1}{|\vec{x} - \vec{y}|^{d+2s-2}} d\vec{y} \right. \\ &\quad \left. - \frac{1}{2s} \int_{\partial K} u(\vec{y}) \frac{\vec{n}_{\vec{y}} \cdot (\vec{x} - \vec{y})}{|\vec{x} - \vec{y}|^{d+2s}} d\vec{y} \right\} \end{aligned}$$

and when $s = 1 - d/2$ then

$$\begin{aligned} \frac{(-\Delta)^s u(\vec{x})}{C(d, s)} &= \int_{\partial K_0} \nabla u|_{K_0} \cdot \vec{n}_{\vec{y}} \log \frac{1}{|\vec{x} - \vec{y}|} d\vec{y} - \frac{u(\vec{x})}{2s} \int_{\partial K_0} \frac{\vec{n}_{\vec{y}} \cdot (\vec{x} - \vec{y})}{|\vec{x} - \vec{y}|^{d+2s}} d\vec{y} \\ &\quad + \sum_{K \neq K_0} \left\{ \frac{1}{2s} \int_{\partial K} \nabla u|_K \cdot \vec{n}_{\vec{y}} \log \frac{1}{|\vec{x} - \vec{y}|} d\vec{y} - \frac{1}{2s} \int_{\partial K} u(\vec{y}) \frac{\vec{n}_{\vec{y}} \cdot (\vec{x} - \vec{y})}{|\vec{x} - \vec{y}|^2} d\vec{y} \right\}. \end{aligned}$$

□

The integrals over the element surface in eqs. (10.15) and (10.16) can be evaluated using standard quadrature rules since the integrals in eqs. (10.15) and (10.16) are regular. Although we are primarily concerned here with the case of piecewise linear basis functions, the above expression can easily be generalized to higher orders finite elements.

Chapter 11

Overcoming Computational Challenges of the Fractional Laplacian

11.1 Computation of Entries of the Stiffness Matrix

The computation of entries of the stiffness matrix A^s in the case of the usual Laplacian ($s = 1$) is straightforward. However, for $s \in (0, 1)$, the bilinear form contains factors $|\vec{x} - \vec{y}|^{-d-2s}$ which means that simple closed forms for the entries are no longer available and suitable quadrature rules therefore must be identified. Moreover, the presence of a repeated integral over Ω (as opposed to an integral over just Ω in the case $s = 1$) means that the matrix needs to be assembled in a double loop over the elements of the mesh so that the computational cost is potentially much larger than in the integer $s = 1$ case. Additionally, every degree of freedom is coupled to all other degrees of freedom and the stiffness matrix is therefore dense.

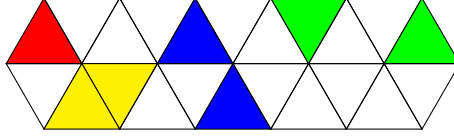


Figure 11.1: Element pairs that are treated separately. We distinguish element pairs of identical elements (*red*), element pairs with common edge (*yellow*), with common vertex (*blue*) and separated elements (*green*).

11.1.1 Reduction to Smooth Integrals

In order to compute the entries of $A^s = \{a(\phi_i, \phi_j)\}_{ij}$ we decompose the expression for the entries into contributions from elements $K, \tilde{K} \in \mathcal{P}_h$ and external edges $e \in \mathcal{P}_{h,\partial}$:

$$a(\phi_i, \phi_j) = \sum_K \sum_{\tilde{K}} a^{K \times \tilde{K}}(\phi_i, \phi_j) + \sum_K \sum_e a^{K \times e}(\phi_i, \phi_j),$$

where the contributions $a^{K \times \tilde{K}}$ and $a^{K \times e}$ are given by:

$$a^{K \times \tilde{K}}(\phi_i, \phi_j) = \frac{C(d, s)}{2} \int_K d\vec{x} \int_{\tilde{K}} d\vec{y} \frac{(\phi_i(\vec{x}) - \phi_i(\vec{y})) (\phi_j(\vec{x}) - \phi_j(\vec{y}))}{|\vec{x} - \vec{y}|^{d+2s}}, \quad (11.1)$$

$$a^{K \times e}(\phi_i, \phi_j) = \frac{C(d, s)}{2s} \int_K d\vec{x} \int_e d\vec{y} \frac{\phi_i(\vec{x}) \phi_j(\vec{x}) \vec{n}_e \cdot (\vec{x} - \vec{y})}{|\vec{x} - \vec{y}|^{d+2s}}. \quad (11.2)$$

Although the following approach holds for arbitrary spatial dimension d , we restrict ourselves to $d = 2$ dimensions. In evaluating the contributions $a^{K \times \tilde{K}}$ over element pairs $K \times \tilde{K}$, several cases need to be distinguished:

1. K and $\tilde{K}$ have empty intersection,
2. K and $\tilde{K}$ are identical,
3. K and $\tilde{K}$ share an edge,
4. K and $\tilde{K}$ share a vertex.

These cases are illustrated in Figure 11.1. In case 1, where the elements do not touch, the Stroud conical quadrature rule [92] (or any other suitable Gauss rule on

simplices) of *sufficiently high order* can be used to approximate the integrals. More details as to what constitutes a sufficiently high order are given in Section 11.1.2.

Special care has to be taken in the remaining cases 2-4, in which the elements are touching, owing to the presence of a singularity in the integrand. The contributions $a^{K \times \tilde{K}}$ and $a^{K \times e}$ as given in eqs. (11.1) and (11.2) for touching elements K and $\tilde{K}$ contain singularities. Fortunately, the singularity is removable and can, as pointed out in [2], be treated using standard techniques from the boundary element literature [82]. More specifically, we write the integral as a sum of integrals over sub-simplices. Each sub-simplex is then mapped onto the hyper-cube $[0, 1]^4$ using the Duffy transformation [44]. The advantage of pursuing this approach is that the singularity arising from the degenerate nature of the Duffy transformation offsets the singularity present in the integrals. Removing the singularity means that the integrals in eqs. (11.1) and (11.2) will become amenable to approximation using standard Gaussian quadrature rules of *sufficiently high order* as discussed in Section 11.1.2. The same idea is applicable in any number of space dimensions.

The derivations of the terms involved can be found in [1, 82] and is outlined below for $d = 2$ dimensions.

The expression for $a^{K \times \tilde{K}}$ can be transformed into integrals over the reference element $\hat{K}$:

$$\begin{aligned} & a^{K \times \tilde{K}}(\phi_i, \phi_j) \\ &= \frac{C(2, s)}{2} \int_K d\vec{x} \int_{\tilde{K}} d\vec{y} \frac{(\phi_i(\vec{x}) - \phi_i(\vec{y})) (\phi_j(\vec{x}) - \phi_j(\vec{y}))}{|\vec{x} - \vec{y}|^{2+2s}} \\ &= \frac{C(2, s)}{2} \frac{|K|}{|\hat{K}|} \frac{|\tilde{K}|}{|\hat{\tilde{K}}|} \int_{\hat{K}} d\hat{\vec{x}} \int_{\hat{\tilde{K}}} d\hat{\vec{y}} \frac{(\phi_i(\vec{x}(\hat{\vec{x}})) - \phi_i(\vec{y}(\hat{\vec{y}}))) (\phi_j(\vec{x}(\hat{\vec{x}})) - \phi_j(\vec{y}(\hat{\vec{y}})))}{|\vec{x}(\hat{\vec{x}}) - \vec{y}(\hat{\vec{y}})|^{2+2s}}. \end{aligned}$$

Similarly, by introducing the reference edge $\hat{e}$, we obtain

$$a^{K \times e}(\phi_i, \phi_j) = \frac{C(2, s)}{2s} \int_K d\vec{x} \int_e d\vec{y} \frac{\phi_i(\vec{x}) \phi_j(\vec{x}) \vec{n}_e \cdot (\vec{x} - \vec{y})}{|\vec{x} - \vec{y}|^{2+2s}}$$

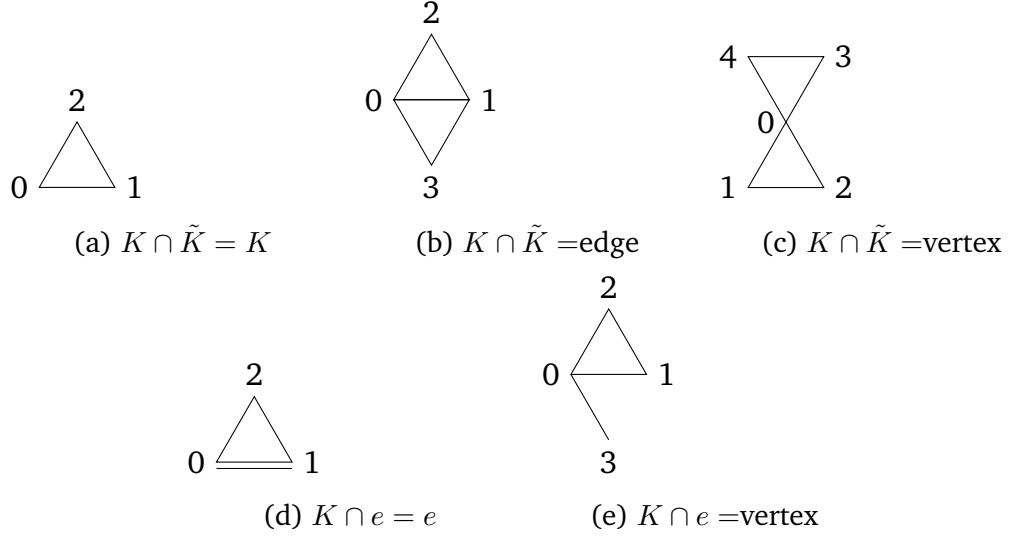


Figure 11.2: Numbering of local nodes for touching triangular elements K and $\tilde{K}$ or element K and edge e .

$$= \frac{C(2, s) |K| |e|}{2^s |\hat{K}| |\hat{e}|} \int_{\hat{K}} d\hat{x} \int_{\hat{e}} d\hat{y} \frac{\phi_i(\vec{x}(\hat{x})) \phi_j(\vec{x}(\hat{x})) \vec{n}_e \cdot (\vec{x}(\hat{x}) - \vec{y}(\hat{y}))}{|\vec{x}(\hat{x}) - \vec{y}(\hat{y})|^{2+2s}}$$

for touching elements K and edges e . If K and $\tilde{K}$ or e have $c \geq 1$ common vertices, and if we designate by λ_k , $k = 0, \dots, 6 - c$ the barycentric coordinates of $K \cup \tilde{K}$ or $K \cup e$ respectively (cf. Figure 11.2), we have

$$\lambda_{k(i)}(\hat{x}) = \phi_i(\vec{x}(\hat{x})),$$

where $k(i)$ is the local index on $K \cup \tilde{K}$ or $K \cup e$ of the global degree of freedom i .

Moreover, we have that

$$\begin{aligned} \vec{x}(\hat{x}) - \vec{y}(\hat{y}) &= \sum_{k=0}^{6-c} \lambda_k(\hat{x}) \vec{x}_k - \sum_{k=0}^{6-c} \lambda_k(\hat{y}) \vec{y}_k \\ &= \sum_{k=0}^{6-c} [\lambda_k(\hat{x}) - \lambda_k(\hat{y})] \vec{x}_k. \end{aligned}$$

Here, $\vec{x}_k$, $k = 0, \dots, 6 - c$ are the vertices that span $K \cup \tilde{K}$ or $K \cup e$ respectively.

By setting

$$\psi_k(\hat{x}, \hat{y}) := \lambda_k(\hat{x}) - \lambda_k(\hat{y}),$$

we can therefore write

$$a^{K \times \tilde{K}}(\phi_i, \phi_j) = \frac{C(2, s)}{2} \frac{|K|}{|\hat{K}|} \frac{|\tilde{K}|}{|\hat{K}|} \int_{\hat{K}} d\hat{x} \int_{\hat{K}} d\hat{y} \frac{\psi_{k(i)}(\hat{x}, \hat{y}) \psi_{k(j)}(\hat{x}, \hat{y})}{\left| \sum_{k=0}^{6-c} \psi_k(\hat{x}, \hat{y}) \vec{x}_k \right|^{2+2s}}.$$

By carefully splitting the integration domain $\hat{K} \times \hat{K}$ into L_c parts and applying a Duffy transformation to each part, the contributions can be rewritten into integrals over a unit hyper-cube, where the singularities are lifted.

$$a^{K \times \tilde{K}}(\phi_i, \phi_j) = \frac{C(2, s)}{2} \frac{|K|}{|\hat{K}|} \frac{|\tilde{K}|}{|\hat{K}|} \sum_{\ell=1}^{L_c} \int_{[0,1]^4} d\vec{\eta} \bar{J}^{(\ell, c)} \frac{\bar{\psi}_{k(i)}^{(\ell, c)}(\vec{\eta}) \bar{\psi}_{k(j)}^{(\ell, c)}(\vec{\eta})}{\left| \sum_{k=0}^{2d-c} \bar{\psi}_k^{(\ell, c)}(\vec{\eta}) \vec{x}_k \right|^{2+2s}}. \quad (11.3)$$

The details of this approach can be found in Chapter 5 of [82] for the interactions between K and $\tilde{K}$. We record the obtained expressions in this case.

- K and $\tilde{K}$ are identical, i.e. $c = 3$

$$L_3 = 3, \quad \bar{J}^{(1,3)} = \bar{J}^{(2,3)} = \bar{J}^{(3,3)} = \eta_0^{3-2s} \eta_1^{2-2s} \eta_2^{1-2s},$$

$$\bar{\psi}_k^{(1,3)} = \begin{cases} -\eta_3 \\ \eta_3 - 1 \\ 1 \end{cases} \quad \bar{\psi}_k^{(2,3)} = \begin{cases} -1 \\ 1 - \eta_3 \\ \eta_3 \end{cases} \quad \bar{\psi}_k^{(3,3)} = \begin{cases} \eta_3 \\ -1 \\ 1 - \eta_3 \end{cases}$$

- K and $\tilde{K}$ share an edge, i.e. $c = 2$

$$L_2 = 5, \quad \bar{J}^{(1,2)} = \eta_0^{3-2s} \eta_1^{2-2s},$$

$$\bar{J}^{(2,2)} = \bar{J}^{(3,2)} = \bar{J}^{(4,2)} = \bar{J}^{(5,2)} = \eta_0^{3-2s} \eta_1^{2-2s} \eta_2$$

$$\bar{\psi}_k^{(1,2)} = \begin{cases} -\eta_2 \\ 1 - \eta_3 \\ \eta_3 \\ \eta_2 - 1 \end{cases} \quad \bar{\psi}_k^{(2,2)} = \begin{cases} -\eta_2 \eta_3 \\ \eta_2 - 1 \\ 1 \\ \eta_2 \eta_3 - \eta_2 \end{cases} \quad \bar{\psi}_k^{(3,2)} = \begin{cases} \eta_2 \\ \eta_2 \eta_3 - 1 \\ 1 - \eta_2 \\ -\eta_2 \eta_3 \end{cases}$$

$$\bar{\psi}_k^{(4,2)} = \begin{cases} \eta_2 \eta_3 \\ 1 - \eta_2 \\ \eta_2 - \eta_2 \eta_3 \\ -1 \end{cases} \quad \bar{\psi}_k^{(5,2)} = \begin{cases} \eta_2 \eta_3 \\ \eta_2 - 1 \\ 1 - \eta_2 \eta_3 \\ -\eta_2 \end{cases}$$

- K and $\tilde{K}$ share a vertex, i.e. $c = 1$

$$L_1 = 2, \quad \bar{J}^{(1,1)} = \bar{J}^{(2,1)} = \eta_0^{3-2s} \eta_2$$

$$\bar{\psi}_k^{(1,1)} = \begin{cases} \eta_2 - 1 \\ 1 - \eta_1 \\ \eta_1 \\ \eta_2 \eta_3 - \eta_2 \\ -\eta_2 \eta_3 \end{cases} \quad \bar{\psi}_k^{(2,1)} = \begin{cases} 1 - \eta_2 \\ \eta_2 - \eta_2 \eta_3 \\ \eta_2 \eta_3 \\ \eta_1 - 1 \\ -\eta_1 \end{cases}$$

We notice that the contributions for identical elements only depend on η_3 , so that in fact only one-dimensional integrals needs to be computed. Similarly, the cases of common edges or common vertices only require two and three dimensional integration.

In a similar fashion, the integration domain of $a^{K \times e}$ can be split into several parts, so that the singularity can be lifted:

$$a^{K \times e}(\phi_i, \phi_j) = \frac{C(2, s)}{2s} \frac{|K|}{|\hat{K}|} \frac{|e|}{|\hat{e}|} \int_{[0,1]^3} d\vec{\eta} \frac{\phi_{k(i)}^{(\ell,c)}(\vec{\eta}) \phi_{k(j)}^{(\ell,c)}(\vec{\eta}) \sum_{k=0}^{5-c} \bar{\psi}_k^{(\ell,c)}(\vec{\eta}) \vec{n}_e \cdot \vec{x}_k}{\left| \sum_{k=0}^{5-c} \bar{\psi}_k^{(\ell,c)}(\vec{\eta}) \vec{x}_k \right|^{2+2s}}. \quad (11.4)$$

Here, $\phi_k^{(\ell,c)}$ are the expressions for the local shape functions under the Duffy transformations. The obtained expressions are

- e is an edge of K , i.e. $c = 2$

$$L_2 = 3, \quad \bar{J}^{(1,2)} = \bar{J}^{(2,2)} = \bar{J}^{(3,2)} = \eta_0^{-2s} (1 - \eta_0),$$

$$\phi_k^{(1,2)} = \begin{cases} 1 - \eta_0 - \eta_2 + \eta_0\eta_2 \\ \eta_0 + \eta_2 - \eta_0\eta_1 - \eta_0\eta_2 \\ \eta_0\eta_1 \end{cases} \quad \phi_k^{(2,2)} = \begin{cases} 1 - \eta_0 - \eta_2 + \eta_0\eta_2 \\ \eta_2 - \eta_0\eta_2 \\ \eta_0 \end{cases}$$

$$\phi_k^{(3,2)} = \begin{cases} 1 - \eta_2 + \eta_0\eta_2 - \eta_0\eta_1 \\ \eta_2 - \eta_0\eta_2 \\ \eta_0\eta_1 \end{cases}$$

$$\bar{\psi}_k^{(1,2)} = \begin{cases} -1 \\ 1 - \eta_1 \\ \eta_1 \end{cases} \quad \bar{\psi}_k^{(2,2)} = \begin{cases} -\eta_1 \\ \eta_1 - 1 \\ 1 \end{cases} \quad \bar{\psi}_k^{(3,2)} = \begin{cases} 1 - \eta_1 \\ -1 \\ \eta_1 \end{cases}$$

We notice that for $s \geq 1/2$, the integrand still contains a singularity. In this case, the finite element space V_h does not include the degrees of freedom on the boundary. For the interaction of the single degree of freedom that is not on the boundary ($k = 2$), we obtain

$$\bar{J}^{(1,2)} = \bar{J}^{(2,2)} = \bar{J}^{(3,2)} = \eta_0^{2-2s} (1 - \eta_0),$$

$$\phi_2^{(1,2)} = \eta_1 \quad \phi_2^{(2,2)} = 1 \quad \phi_2^{(3,2)} = \eta_1$$

and $\bar{\psi}_2^{\ell,c}$ as above.

- K and e share a vertex, i.e. $c = 1$

$$L_1 = 2, \quad \bar{J}^{(1,1)} = \eta_0^{1-2s}, \bar{J}^{(2,1)} = \eta_0^{1-2s}\eta_1$$

$$\bar{\psi}_k^{(1,1)} = \begin{cases} \eta_2 - 1 \\ 1 - \eta_1 \\ \eta_1 \\ -\eta_2 \end{cases} \quad \bar{\psi}_k^{(2,1)} = \begin{cases} 1 - \eta_1 \\ \eta_1 - \eta_1\eta_2 \\ \eta_1\eta_2 \\ -1 \end{cases}$$

11.1.2 Determining the Order of the Quadrature Rules

The foregoing considerations show that the evaluation of the entries of the stiffness matrix boils down to the evaluation of integrals with smooth integrands, i.e. expressions eqs. (11.1) and (11.2) for case 1 and expressions eqs. (11.3) and (11.4) for case 2-4. As mentioned earlier, it is necessary to use a *sufficiently high order* quadrature rule to approximate these integrals. We now turn to the question of how high is sufficient.

The arguments used to prove the ensuing estimates follow a pattern similar to the proofs of Theorems 5.3.29, 5.3.23 and 5.3.24 in [82]. The main difference from [82] is the presence of the boundary integral term.

Theorem 39 ([82], Theorems 5.3.23 and 5.3.24).

If K and $\tilde{K}$ (K and e) are touching elements, then

$$\begin{aligned} |E_{K \times \tilde{K}}^{i,j}| &\leq C h_K^{2-2s} \rho_1^{-2k_T}, \\ |E_{K \times e}^{i,j}| &\leq C h_K^{2-2s} \rho_3^{-2k_{T,\partial}}, \end{aligned}$$

where $\rho_1, \rho_3 > 1$ and $k_T, k_{T,\partial}$ are the quadrature orders in every dimension of eqs. (11.1) and (11.2) (after regularization).

If K and $\tilde{K}$ (K and e) are not touching, then

$$|E_{K \times \tilde{K}}^{i,j}| \leq C (h_K h_{\tilde{K}})^2 d_{K,\tilde{K}}^{-2-2s} \left\{ \left(\frac{d_{K,\tilde{K}}}{h_K} \right)^2 \tilde{\rho}_2 (K, \tilde{K})^{-2k_{NT}} + \left(\frac{d_{K,\tilde{K}}}{h_{\tilde{K}}} \right)^2 \tilde{\rho}_2 (\tilde{K}, K)^{-2k_{NT}} \right\},$$

$$|E_{K \times e}^{i,j}| \leq C (h_K h_e)^2 d_{K,e}^{-2-2s} \left\{ \left(\frac{d_{K,e}}{h_K} \right)^2 \tilde{\rho}_4(K, e)^{-2k_{NT}} + \left(\frac{d_{K,e}}{h_e} \right)^2 \tilde{\rho}_4(e, K)^{-2k_{NT}} \right\},$$

where $d_{K,\tilde{K}} := \text{dist}(K, \tilde{K})$, $d_{K,e} := \text{dist}(K, e)$, $\tilde{\rho}_2(K, \tilde{K}) := \rho_2 \max \left\{ \frac{d_{K,\tilde{K}}}{h_K}, 1 \right\}$, $\tilde{\rho}_4(K, e) := \rho_4 \max \left\{ \frac{d_{K,e}}{h_K}, 1 \right\}$ and $\rho_2, \rho_4 > 1$, and k_{NT} , $k_{NT,\partial}$ are the quadrature order in every dimension of eqs. (11.1) and (11.2).

Theorem 40.

For $d = 2$, let $\mathcal{I}_K$ index the degrees of freedom on $K \in \mathcal{P}_h$, and define $\mathcal{I}_{K \times \tilde{K}} := \mathcal{I}_K \cup \mathcal{I}_{\tilde{K}}$. Let $k_T(K, \tilde{K})$ (respectively $k_{T,\partial}(K, e)$) be the quadrature order used for touching pairs $K \times \tilde{K}$ (respectively $K \times e$), and let $k_{NT}(K, \tilde{K})$ (respectively $k_{NT,\partial}(K, e)$) be the quadrature order used for pairs that have empty intersection. Denote the resulting approximation to the bilinear form $a(\cdot, \cdot)$ by $a^Q(\cdot, \cdot)$. Then the consistency error due to quadrature is bounded by

$$|a(u, v) - a^Q(u, v)| \leq C (E_T + E_{NT} + E_{T,\partial} + E_{NT,\partial}) \|u\|_{L^2(\Omega)} \|v\|_{L^2(\Omega)} \quad \forall u, v \in V_h,$$

where the errors are given by

$$\begin{aligned} E_T &= n \max_{K, \tilde{K} \in \mathcal{P}_h, \overline{K} \cap \overline{\tilde{K}} \neq \emptyset} h_K^{-2s} \rho_1^{-2k_T(K, \tilde{K})}, \\ E_{NT} &= n \max_{K, \tilde{K} \in \mathcal{P}_h, \overline{K} \cap \overline{\tilde{K}} = \emptyset} \left[h_{\tilde{K}}^2 \tilde{\rho}_2(K, \tilde{K})^{-2k_{NT}(K, \tilde{K})} + h_K^2 \tilde{\rho}_2(\tilde{K}, K)^{-2k_{NT}(K, \tilde{K})} \right] \\ &\quad d_{K,\tilde{K}}^{-2s} \max \left\{ h_K^{-2}, h_{\tilde{K}}^{-2} \right\}, \\ E_{T,\partial} &= n^{1/2} \max_{K \in \mathcal{P}_h, e \in \partial \mathcal{P}_h, \overline{K} \cap \overline{e} \neq \emptyset} h_K^{-2s} \rho_3^{-2k_{T,\partial}(K, e)}, \\ E_{NT,\partial} &= n^{1/2} \max_{K \in \mathcal{P}_h, e \in \partial \mathcal{P}_h, \overline{K} \cap \overline{e} = \emptyset} \left[h_e^2 \tilde{\rho}_4(K, e)^{-2k_{NT}(K, e)} + h_K^2 \tilde{\rho}_4(e, K)^{-2k_{NT}(K, e)} \right] d_{K,e}^{-2s} h_K^{-2}, \end{aligned}$$

and where $d_{K,\tilde{K}} := \inf_{\vec{x} \in K, \vec{y} \in \tilde{K}} |\vec{x} - \vec{y}|$, $d_{K,e} := \inf_{\vec{x} \in K, \vec{y} \in e} |\vec{x} - \vec{y}|$, and $\rho_1, \tilde{\rho}_2, \rho_3, \tilde{\rho}_4$ are as in Theorem 39.

Proof of Theorem 40. Let the quadrature rules for the pairs $K \times \tilde{K}$ and $K \times e$ be denoted by $a_Q^{K \times \tilde{K}}(\cdot, \cdot)$ and $a_Q^{K \times e}(\cdot, \cdot)$. Set

$$E_{K \times \tilde{K}}^{i,j} = a^{K \times \tilde{K}}(\phi_i, \phi_j) - a_Q^{K \times \tilde{K}}(\phi_i, \phi_j),$$

$$E_{K \times e}^{i,j} = a^{K \times e}(\phi_i, \phi_j) - a_Q^{K \times e}(\phi_i, \phi_j).$$

For $u, v \in V_h$, we set

$$\begin{aligned} E_{K \times \tilde{K}}(u, v) &= \sum_{i \in \mathcal{I}_{K \times \tilde{K}}} \sum_{j \in \mathcal{I}_{K \times \tilde{K}}} u_i v_j E_{K \times \tilde{K}}^{i,j}, \\ E_{K \times e}(u, v) &= \sum_{i \in \mathcal{I}_K} \sum_{j \in \mathcal{I}_K} u_i v_j E_{K \times e}^{i,j} \end{aligned}$$

so that

$$\begin{aligned} |E_{K \times \tilde{K}}(u, v)| &\leq \left(\max_{i,j} |E_{K \times \tilde{K}}^{i,j}| \right) \sum_{i \in \mathcal{I}_{K \times \tilde{K}}} |u_i| \sum_{j \in \mathcal{I}_{K \times \tilde{K}}} |v_j| \\ &\leq \left(\max_{i,j} |E_{K \times \tilde{K}}^{i,j}| \right) |\mathcal{I}_{K \times \tilde{K}}| \sqrt{\sum_{i \in \mathcal{I}_{K \times \tilde{K}}} |u_i|^2} \sqrt{\sum_{j \in \mathcal{I}_{K \times \tilde{K}}} |v_j|^2}, \\ |E_{K \times e}(u, v)| &\leq \left(\max_{i,j} |E_{K \times e}^{i,j}| \right) \sum_{i \in \mathcal{I}_K} |u_i| \sum_{j \in \mathcal{I}_K} |v_j| \\ &\leq \left(\max_{i,j} |E_{K \times e}^{i,j}| \right) |\mathcal{I}_K| \sqrt{\sum_{i \in \mathcal{I}_K} |u_i|^2} \sqrt{\sum_{j \in \mathcal{I}_K} |v_j|^2} \end{aligned}$$

Since

$$\begin{aligned} \sum_{i \in \mathcal{I}_{K \times \tilde{K}}} |u_i|^2 &\leq C \left[h_K^{-d} \int_K u^2 + h_{\tilde{K}}^{-d} \int_{\tilde{K}} u^2 \right], \\ \sum_{i \in \mathcal{I}_K} |u_i|^2 &\leq C h_K^{-d} \int_K u^2, \end{aligned}$$

we find

$$\begin{aligned} |a(u, v) - a_Q(u, v)| &\leq \sum_K \sum_{\tilde{K}} |E_{K \times \tilde{K}}(u, v)| + \sum_K \sum_e |E_{K \times e}(u, v)| \\ &\leq C \sum_K \sum_{\tilde{K}} \left(\max_{i,j} |E_{K \times \tilde{K}}^{i,j}| \right) \max \{ h_K^{-d}, h_{\tilde{K}}^{-d} \} \left[\|u\|_{L^2(K)}^2 + \|u\|_{L^2(\tilde{K})}^2 \right]^{1/2} \\ &\quad \left[\|v\|_{L^2(K)}^2 + \|v\|_{L^2(\tilde{K})}^2 \right]^{1/2} \\ &\quad + C \sum_K \sum_e \left(\max_{i,j} |E_{K \times e}^{i,j}| \right) h_K^{-d} \|u\|_{L^2(K)} \|v\|_{L^2(K)} \end{aligned}$$

$$\begin{aligned}
&\leq C \left(\max_{K, \tilde{K}} \max_{i,j} |E_{K \times \tilde{K}}^{i,j}| \max \{h_K^{-d}, h_{\tilde{K}}^{-d}\} \right) \sum_K \sum_{\tilde{K}} \|u\|_{L^2(K \cup \tilde{K})} \|v\|_{L^2(K \cup \tilde{K})} \\
&\quad + C \left(\max_{K,e} \max_{i,j} |E_{K \times e}^{i,j}| h_K^{-d} \right) \sum_K \sum_e \|u\|_{L^2(K)} \|v\|_{L^2(K)}.
\end{aligned}$$

Because

$$\begin{aligned}
\sum_K \sum_{\tilde{K}} \|u\|_{L^2(K \cup \tilde{K})} \|v\|_{L^2(K \cup \tilde{K})} &\leq \sqrt{\sum_K \sum_{\tilde{K}} \|u\|_{L^2(K \cup \tilde{K})}^2} \sqrt{\sum_K \sum_{\tilde{K}} \|v\|_{L^2(K \cup \tilde{K})}^2} \\
&\leq 2 |\mathcal{P}_h| \|u\|_{L^2(\Omega)} \|v\|_{L^2(\Omega)} \\
&\leq Cn \|u\|_{L^2(\Omega)} \|v\|_{L^2(\Omega)}
\end{aligned}$$

and

$$\begin{aligned}
\sum_K \sum_e \|u\|_{L^2(K)} \|v\|_{L^2(K)} &\leq |\partial \mathcal{P}_h| \|u\|_{L^2(\Omega)} \|v\|_{L^2(\Omega)} \\
&\leq Cn^{(d-1)/d} \|u\|_{L^2(\Omega)} \|v\|_{L^2(\Omega)},
\end{aligned}$$

we obtain

$$\begin{aligned}
|a(u, v) - a_Q(u, v)| &\leq C \left[n \left(\max_{K, \tilde{K}} \max_{i,j} |E_{K \times \tilde{K}}^{i,j}| \max \{h_K^{-d}, h_{\tilde{K}}^{-d}\} \right) \right. \\
&\quad \left. + n^{(d-1)/d} \left(\max_{K,e} \max_{i,j} |E_{K \times e}^{i,j}| h_K^{-d} \right) \right] \|u\|_{L^2(\Omega)} \|v\|_{L^2(\Omega)}.
\end{aligned}$$

For $d = 2$, using Theorem 39 stated below permits to conclude. $\square$

The impact of the use of quadrature rules on the accuracy of the resulting finite element approximation can be quantified using Strang's first lemma [50, Lemma 2.27]:

$$\begin{aligned}
\|u - u_h\|_{\tilde{H}^s(\Omega)} &\leq C \inf_{v_h \in V_h} \left[\|u - v_h\|_{\tilde{H}^s(\Omega)} + \sup_{w_h \in V_h} \frac{|a(v_h, w_h) - a_Q(v_h, w_h)|}{\|w_h\|_{\tilde{H}^s(\Omega)}} \right] \\
&\leq C \inf_{v_h \in V_h} \left[\|u - v_h\|_{\tilde{H}^s(\Omega)} \right. \\
&\quad \left. + (E_T + E_{NT} + E_{T,\partial} + E_{NT,\partial}) \|v_h\|_{L^2(\Omega)} \sup_{w_h \in V_h} \frac{\|w_h\|_{L^2(\Omega)}}{\|w_h\|_{\tilde{H}^s(\Omega)}} \right]
\end{aligned}$$

$$\leq C \inf_{v_h \in V_h} \left[\|u - v_h\|_{\tilde{H}^s(\Omega)} + (E_T + E_{NT} + E_{T,\partial} + E_{NT,\partial}) \|v_h\|_{L^2(\Omega)} \right],$$

where we used the Poincare inequality $\|w_h\|_{L^2(\Omega)} \leq C \|w_h\|_{\tilde{H}^s(\Omega)}$ in the last step. Based on the error estimates of Theorem 40, the optimal quadrature orders in order to allow convergence of up to order $n^{-\alpha}$ should be chosen as

$$\begin{aligned} & k_T(K, \tilde{K}) \\ & \geq \frac{(\alpha/2 + 1/2) \log n + s |\log h_K|}{\log(\rho_1)} - C, \\ & k_{NT}(K, \tilde{K}) \\ & \geq \max \left\{ \frac{(\alpha/2 + 1/2) \log n + (s-1) |\log h_{\tilde{K}}| + \max\{|\log h_K|, |\log h_{\tilde{K}}|\} - s \log \frac{d_{K,\tilde{K}}}{h_{\tilde{K}}} - C}{\log_+ \frac{d_{K,\tilde{K}}}{h_K} + \log(\rho_2)}, \right. \\ & \quad \left. \frac{(\alpha/2 + 1/2) \log n + (s-1) |\log h_K| + \max\{|\log h_K|, |\log h_{\tilde{K}}|\} - s \log \frac{d_{K,\tilde{K}}}{h_K} - C}{\log_+ \frac{d_{K,\tilde{K}}}{h_{\tilde{K}}} + \log(\rho_2)} \right\} \\ & k_{T,\partial}(K, e) \\ & \geq \frac{(\alpha/2 + 1/4) \log n + s |\log h_K|}{\log(\rho_3)} - C, \\ & k_{NT,\partial}(K, e) \\ & \geq \max \left\{ \frac{(\alpha/2 + 1/4) \log n + (s-1) |\log h_e| + \max\{|\log h_K|, |\log h_e|\} - s \log \frac{d_{K,e}}{h_e} - C}{\log_+ \frac{d_{K,e}}{h_K} + \log(\rho_4)}, \right. \\ & \quad \left. \frac{(\alpha/2 + 1/4) \log n + (s-1) |\log h_K| + \max\{|\log h_K|, |\log h_e|\} - s \log \frac{d_{K,e}}{h_K} - C}{\log_+ \frac{d_{K,e}}{h_e} + \log(\rho_4)} \right\}. \end{aligned}$$

It transpires from the expressions derived above and the fact that $n \sim h^{-2}$ for a quasi-uniform mesh that the complexity to calculate the contributions by a single pair of elements K and $\tilde{K}$ scales like

- $\log n$ if the elements coincide,
- $(\log n)^2$ if the elements share only an edge,
- $(\log n)^3$ if the elements share only a vertex,

- $(\log n)^4$ if the elements have empty intersection, but are “near neighbors”, and
- C if the elements are well separated.

Since $n \sim |\mathcal{P}_h|$, we cannot expect a straightforward assembly of the stiffness matrix to scale better than $\mathcal{O}(n^2)$. Similarly, its memory requirement is n^2 , and a single matrix-vector product has complexity $\mathcal{O}(n^2)$, which severely limits the size of problems that can be considered.

11.2 Solving the Linear Systems

The fractional Poisson equation (10.1) leads to a dense linear algebraic system

$$\mathbf{A}^s \vec{u} = \vec{b}. \quad (11.5)$$

The solution of the dense matrix equation (11.5) also poses its own problems. For instance, using matrix factorization to solve the resulting system has complexity $\mathcal{O}(n^3)$, where n is the number of degrees of freedom.

Alternatively, if conjugate gradient iteration is used, the number of iterations necessary to converge to a given error tolerance is known to scale as the square root of the condition number of the matrix. The condition numbers of the fractional Laplacian, as well as the its diagonally preconditioned version grow with the number of unknowns n , as shown by Theorem 41.

Theorem 41 ([10]).

For $s < d/2$, and a family of shape regular triangulations $\mathcal{P}_h$ with minimal and maximal element size $h_{\min}$ and h , the spectrum of the stiffness matrix satisfies

$$\begin{aligned} cn^{-2s/d} h_{\min}^{d-2s} \mathbf{I} &\leq \mathbf{A}^s \leq Ch^{d-2s} \mathbf{I}, \\ cn^{-2s/d} \mathbf{I} &\leq (\mathbf{D}^s)^{-1/2} \mathbf{A}^s (\mathbf{D}^s)^{-1/2} \leq C \mathbf{I}, \end{aligned}$$

where $\mathbf{D}^s$ is the diagonal part of $\mathbf{A}^s$. The condition number of the stiffness matrix satisfies

$$\kappa(\mathbf{A}^s) = C \left(\frac{h}{h_{\min}} \right)^{d-2s} n^{2s/d}, \quad (11.6)$$

$$\kappa((\mathbf{D}^s)^{-1} \mathbf{A}^s) = C n^{2s/d}. \quad (11.7)$$

The exponent of the growth of the condition number depends on the fractional order s . For small s , the matrix is better conditioned. For large s , the growth of the condition number approaches that of the usual integer order Laplacian. Moreover, (11.6) shows that the ill-conditioning becomes increasingly severe on non-uniform meshes. Fortunately, estimate (11.7) shows that using diagonal scaling removes the effects of the non-uniformity of the triangulation. The conjugate gradient method applied to eq. (11.5) is therefore expected to converge in $\sqrt{\kappa((\mathbf{D}^s)^{-1} \mathbf{A}^s)} \sim n^{s/d}$ iterations. Alternatively, a geometric multigrid solver can be used. In the integer order case, multigrid iterations are used to good effect for solving systems involving both the mass matrix and the stiffness matrix which arises from the discretisation of the regular Laplacian. It is known that multigrid converges with a rate that is independent of the problem size for pseudo-differential operators of positive order in general, which in particular includes the fractional Laplacian [62, 63, 82]. Although a single multigrid iteration is more expensive than a single iteration of conjugate gradient, a multigrid solver is more efficient, especially for larger fractional order s that are troublesome for the conjugate gradient method.

As mentioned earlier, the temporal problem (10.2) is often of more interest than the steady state problem (10.1). We first observe that using an explicit schemes for the time discretization will lead to a CFL conditions on the time-step size Δt of type $\Delta t \leq C h^{2s}$. The discretization using an implicit integration scheme in the time variable lead to systems of the form

$$(\mathbf{M} + \Delta t \mathbf{A}^s) \vec{u} = \vec{b}, \quad (11.8)$$

where $\mathbf{M}$ is the mass matrix. In typical examples, the time-step will be chosen so that the error from the approximation in time matches the spatial order, and hence Δt depends on the spatial discretisation. The following Theorem shows that if the time-step size is $\mathcal{O}(h^{2s})$, we can expect the conjugate gradient method to converge rapidly, since the condition number of the system is bounded by a constant.

Lemma 42.

For a shape regular and globally quasi-uniform family of triangulations $\mathcal{P}_h$ and time-step $\Delta t \leq 1$,

$$\kappa(\mathbf{M} + \Delta t \mathbf{A}^s) \leq C \left(1 + \frac{\Delta t}{h^{2s}}\right).$$

More generally, for a family of triangulations that is only locally quasi-uniform and $\Delta t \leq h_{\min}^{2s} n^{2s/d}$, it holds that

$$\kappa(\mathbf{M} + \Delta t \mathbf{A}^s) \leq C \left(\frac{h}{h_{\min}}\right)^d \left(1 + \frac{\Delta t}{h^{2s}}\right). \quad (11.9)$$

If $\mathbf{D}^0$ is taken to be the diagonal part of the mass matrix and $\Delta t \leq h^{2s} n^{2s/d}$, then

$$\kappa\left((\mathbf{D}^0)^{-1}(\mathbf{M} + \Delta t \mathbf{A}^s)\right) \leq C \left(1 + \frac{\Delta t}{h_{\min}^{2s}}\right) \quad (11.10)$$

Proof. Since $ch_{\min}^d \mathbf{I} \leq \mathbf{M} \leq Ch^d \mathbf{I}$, this also permits us to deduce that

$$c(h_{\min}^d + \Delta t n^{-2s/d} h_{\min}^{d-2s}) \mathbf{I} \leq \mathbf{M} + \Delta t \mathbf{A}^s \leq C(h^d + \Delta t h^{d-2s}) \mathbf{I}$$

and so

$$\begin{aligned} \kappa(\mathbf{M} + \Delta t \mathbf{A}^s) &\leq C \frac{h^d + \Delta t h^{d-2s}}{h_{\min}^d + \Delta t n^{-2s/d} h_{\min}^{d-2s}} \\ &= C \left(\frac{h}{h_{\min}}\right)^d \frac{1 + \Delta t h^{-2s}}{1 + \Delta t n^{-2s/d} h_{\min}^{-2s}} \\ &\leq C \left(\frac{h}{h_{\min}}\right)^d \left(1 + \frac{\Delta t}{h^{2s}}\right). \end{aligned}$$

For globally quasi-uniform meshes, $\frac{h}{h_{\min}} \sim C$ and $h_{\min}^{2s} n^{2s/d} \sim 1$.

Since $\|\phi_i\|_{L^2}^2 \sim h_i^d$ and $\|\phi_i\|_{\tilde{H}^s}^2 \sim h_i^{d-2s}$ [10] and $(D^0)^{-1/2} M (D^0)^{-1/2} \sim I$, we find

$$c(1 + \Delta t n^{-2s/d} h^{-2s}) I \leq (D^0)^{-1/2} (M + \Delta t A^s) (D^0)^{-1/2} \leq C(1 + \Delta t h_{\min}^{-2s}) I,$$

and hence

$$\kappa \left((D^0)^{-1} (M + \Delta t A^s) \right) \leq C \frac{1 + \Delta t h_{\min}^{-2s}}{1 + \Delta t n^{-2s/d} h^{-2s}} \leq C \left(1 + \frac{\Delta t}{h_{\min}^{2s}} \right).$$

□

Lemma 42 shows that if the time step Δt is chosen to be $\Delta t \sim h_{\min}^{2s}$, then using the conjugate gradient method to solve the diagonally preconditioned system (11.8) gives a uniform bound on the number of iterations. However, if larger time steps are chosen, e.g. independent of $h_{\min}$, then (11.10) shows that the simple diagonal scaling becomes inefficient resulting in a number of iterations growing as $\mathcal{O}\left(\frac{\Delta t}{h_{\min}^{2s}}\right)$. Nevertheless, by applying a multigrid solver (as in the steady state case), one can restore a uniform bound on the number of iterations [62, 63, 82].

In this section we have concerned ourselves with the effect the mesh and the fractional order has on the rate of convergence of iterative solvers. This, of course, ignores the cost of carrying out the iteration. The complexity of both multigrid and conjugate gradient iterations depends on how efficiently the matrix-vector product $A^s \vec{x}$ can be computed. (The mass matrix in eq. (11.8) has $\mathcal{O}(n)$ entries, so its matrix-vector product scales linearly in the number of unknowns.) Since all the basis functions ϕ_i interact with each other, the matrix A^s is dense and its matrix-vector product has complexity n^2 . In the following section, we discuss a sparse approximation that will preserve the order of the approximation error of the fractional Laplacian, but display significantly better scaling in memory and operations for both assembly and matrix-vector product.

11.3 Sparse Approximation of the Stiffness Matrix and of the Strong Form

The presence of a factor $|\vec{x} - \vec{y}|^{-d-2s}$ in the integrand in the expression for the entries of the stiffness matrix means that the contribution of pairs of elements that are well separated is significantly smaller than the contribution arising from pairs of elements that are close to one another. The evaluation of the strong form of the fractional Laplacian as given in eqs. (10.15) and (10.16) involves integration over all edges of the mesh weighted by the same kernel. This suggests the use of the *cluster method* from the boundary element literature, whereby such far field contributions are replaced by less expensive low-rank blocks rather than computing and storing the all individual entries from the original matrix. Conversely, the near-field contributions are more significant but involve only local couplings and hence the cost of storing the individual entries is a practical proposition. A full discussion of the clustering method is beyond the scope of the present work but can be found in [82, Chapter 7]. Here, we confine ourselves to stating only the necessary definitions and steps needed to describe our approach.

Definition 43 ([82]).

A cluster is a union of one or more indices from the set of degrees of freedom $\mathcal{I}$. The nodes of a hierarchical cluster tree $\mathcal{T}$ are clusters. The set of all nodes is denoted by T and satisfies

1. $\mathcal{I}$ is a node of $\mathcal{T}$.
2. The set of leaves $Leaves(\mathcal{T}) \subset T$ corresponds to the degrees of freedom $i \in \mathcal{I}$ and is given by

$$Leaves(\mathcal{T}) := \{\{i\} : i \in \mathcal{I}\}.$$

3. For every $\sigma \in T \setminus Leaves(\mathcal{T})$ there exists a minimal set $\Sigma(\sigma)$ of nodes in $T \setminus \{\sigma\}$

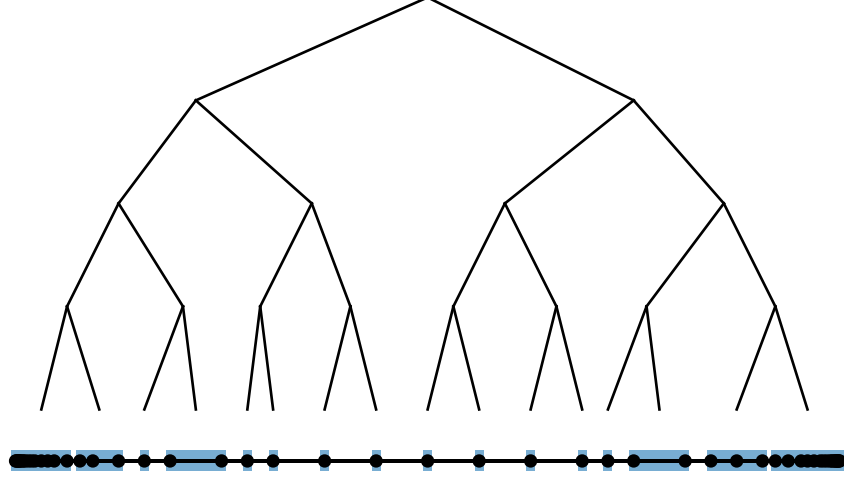


Figure 11.3: Cluster tree for a one-dimensional problem. The mesh with its nodal degrees of freedom is plotted at the bottom. The leaf clusters are colored in blue.

that satisfies

$$\sigma = \bigcup_{\tau \in \Sigma(\sigma)} \tau.$$

The set $\Sigma(\sigma)$ is called the sons of σ . The edges of the cluster tree $\mathcal{T}$ are the pairs of nodes $(\sigma, \tau) \in T \times T$ such that $\tau \in \Sigma(\sigma)$.

An example of a cluster tree for a one-dimensional problem is given in Figure 11.3.

Definition 44 ([82]).

The cluster box Q_σ of a cluster $\sigma \in T$ is the minimal hyper-cube which contains $\bigcup_{i \in \sigma} \text{supp } \phi_i$. The diameter of a cluster is the diameter of its cluster box $\text{diam}(\sigma) := \sup_{\vec{x}, \vec{y} \in Q_\sigma} |\vec{x} - \vec{y}|$. The distance of two clusters σ and τ is $\text{dist}(\sigma, \tau) := \inf_{\vec{x} \in Q_\sigma, \vec{y} \in Q_\tau} |\vec{x} - \vec{y}|$. The subspace V_σ of V_h is defined as $V_\sigma := \text{span} \{\phi_i \mid i \in \sigma\}$.

For given $\eta > 0$, a pair of clusters (σ, τ) is called *admissible*, if

$$\eta \text{dist}(\sigma, \tau) \geq \max \{\text{diam}(\sigma), \text{diam}(\tau)\}.$$

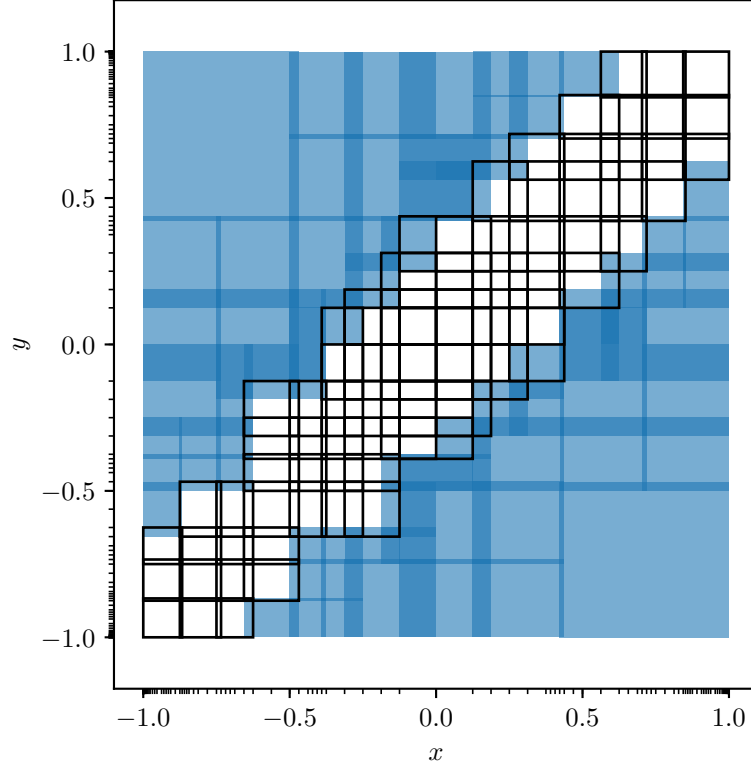


Figure 11.4: Cluster pairs for a one-dimensional problem. The cluster boxes of the admissible cluster pairs are colored in light blue, and their overlap in darker blue. The diagonal cluster pairs are not admissible and are not approximated, but assembled in full. The nodes of the nodal degrees of freedom are plotted on x - and y -axis.

The admissible cluster pairs can be determined recursively. Cluster pairs that are not admissible and have no admissible sons are part of the near field and are assembled into a sparse matrix. The admissible cluster pairs for a one dimensional problem are shown in Figure 11.4.

For admissible pairs of clusters σ and τ and any functions ϕ, ψ with $\text{supp } \phi \subset Q_\sigma$ and $\text{supp } \psi \subset Q_\tau$, the bilinear form a evaluates to

$$a(\phi, \psi) = -C(d, s) \int_{\Omega} \int_{\Omega} k(\vec{x}, \vec{y}) \phi(\vec{x}) \psi(\vec{y}),$$

since $Q_\sigma \cap Q_\tau = \emptyset$. Here, $k(\vec{x}, \vec{y}) = |\vec{x} - \vec{y}|^{-(d+2s)}$ is the interaction kernel, which can be approximated on $Q_\sigma \times Q_\tau$ using Chebyshev interpolation of order m in every

spatial dimension by

$$k_m(\vec{x}, \vec{y}) = \sum_{\alpha, \beta=1}^{m^d} k(\vec{\xi}_\alpha^\sigma, \vec{\xi}_\beta^\tau) L_\alpha^\sigma(\vec{x}) L_\beta^\tau(\vec{y}).$$

L_α^σ are the tensor Chebyshev polynomials of order m on the cluster box Q_σ , and $\vec{\xi}_\alpha^\sigma$ are the tensor Chebyshev points on Q_σ with $L_\alpha^\sigma(\vec{\xi}_\beta^\sigma) = \delta_{\alpha\beta}$. This leads to the following approximation:

$$a(\phi, \psi) \approx -C(d, s) \sum_{\alpha, \beta=1}^{m^2} k(\vec{\xi}_\alpha^\sigma, \vec{\xi}_\beta^\tau) \int_{\text{supp } \phi} \phi(\vec{x}) L_\alpha^\sigma(\vec{x}) d\vec{x} \int_{\text{supp } \psi} \psi(\vec{y}) L_\beta^\tau(\vec{y}) d\vec{y}.$$

Expressions of the type $\int_{\text{supp } \phi} \phi(\vec{x}) L_\alpha^\sigma(\vec{x}) d\vec{x}$ can be computed recursively starting from the finest level of the cluster tree, since for $\tau \in \Sigma(\sigma)$ and $\vec{x} \in Q_\tau$

$$L_\alpha^\sigma(\vec{x}) = \sum_{\beta} L_\alpha^\sigma(\vec{\xi}_\beta^\tau) L_\beta^\tau(\vec{x}).$$

This means that for all leaves $\sigma = \{i\}$, and all $1 \leq \alpha \leq m^d$, the *basis far-field coefficients*

$$\int_{\text{supp } \phi \cap Q_\sigma} \phi(\vec{x}) L_\alpha^\sigma(\vec{x}) d\vec{x}$$

need to be evaluated (e.g. by $m + 1$ -th order Gaussian quadrature). Moreover, the *shift coefficients*

$$L_\alpha^\sigma(\vec{\xi}_\beta^\tau)$$

for $\tau \in \Sigma(\sigma)$ must be computed, as well as the kernel approximations

$$k(\vec{\xi}_\alpha^\sigma, \vec{\xi}_\beta^\tau)$$

for every admissible pair of clusters (σ, τ) . We refer the reader to [82] for further details.

In order for the consistency error of the cluster method to be dominated by the discretization error of the method, the interpolation order m essentially needs to grow with $|\log h_{\min}|$. For further details, we refer to [9, 82].

By following the arguments in [82], it can be shown that the number of near field entries scales linearly in n . The same conclusion holds for the number of far field cluster pairs. Since the four dimensional integral contributions $a^{K \times \tilde{K}}$ are evaluated using Gaussian quadrature rules with at most $k \sim \log n$ quadrature nodes per dimension, the assembly of the near field contributions scales with $n \log^{2d} n$. The far field kernel approximations and the shift coefficients have size $m^{2d} \sim \log^{2d} n$, and are also calculated in $\log^{2d} n$ complexity. This means that all the kernel approximations and shift coefficients are obtained in $n \log^{2d} n$ time. Finally, the nm^d basis far-field coefficients require the evaluation of integrals using $m+1$ -th order Gaussian quadrature, leading to a complexity of $n \log^{2d} n$ as well. The overall complexity of the cluster method is therefore $\mathcal{O}(n \log^{2d} n)$, and the sparse approximation requires $\mathcal{O}(n \log^{2d} n)$ memory. In practice, this means that the assembly of the near-field matrix dominates the other steps but involves only local computations.

The computation of the matrix-vector product involving upward and downward recursion in the cluster tree and multiplication by the kernel approximations can also be shown to scale with $\mathcal{O}(n \log^{2d} n)$.

The same principle can be used to evaluate the strong form in all the quadrature nodes necessary for the computation of the BR error indicators. According to eqs. (10.15) and (10.16), a naive implementation scales as $\mathcal{O}(n^2)$. Many of the components that are used for the cluster method in the matrix-vector product can be reused. Let $u = \sum_{i \in \mathcal{I}_h} u_i \phi_i \in V_h$, and let $\vec{x}_0 \in K_0$ be a quadrature node in the evaluation of the local L^2 -norms in (10.13). Moreover, let $i_0 \in \mathcal{I}_h$ be the closest degree of freedom to $\vec{x}_0$ on K_0 .

Let (σ, τ) be an admissible cluster pair such that $i_0 \in \sigma$, and set $u_\tau = \sum_{j \in \tau} u_j \phi_j$, so that $\text{supp } u_\tau \subset Q_\tau$. The contribution to the fractional Laplacian is approximated by

$$- \int_{Q_\tau} k(\vec{x}, \vec{y}) u_\tau(\vec{y}) d\vec{y} = - \int_{Q_\sigma} \int_{Q_\tau} k(\vec{x}, \vec{y}) \delta(\vec{x} - \vec{x}_0) u_\tau(\vec{y}) d\vec{y} d\vec{x}$$

$$\begin{aligned}
&\approx \sum_{\alpha, \beta=1}^{m^d} k \left(\vec{\xi}_\alpha^\sigma, \vec{\xi}_\beta^\tau \right) \int_{Q_\sigma} L_\alpha^\sigma(\vec{x}) \delta(\vec{x} - \vec{x}_0) d\vec{x} \int_{Q_\tau} L_\beta^\tau(\vec{y}) u_\tau(\vec{y}) d\vec{y} \\
&= \sum_{\alpha, \beta=1}^{m^d} k \left(\vec{\xi}_\alpha^\sigma, \vec{\xi}_\beta^\tau \right) L_\alpha^\sigma(\vec{x}_0) \int_{Q_\tau} L_\beta^\tau(\vec{y}) u_\tau(\vec{y}) d\vec{y}.
\end{aligned}$$

The integrals $\int_{Q_\tau} L_\beta^\tau(\vec{y}) u_\tau(\vec{y}) d\vec{y}$ already appeared in the computation of the matrix-vector product, and can be recursively computed. The same holds for the factors $L_\alpha^\sigma(\vec{x}_0)$.

On the other hand, if the cluster pair (σ, τ) is not admissible and therefore in the near-field, then formulas eqs. (10.15) and (10.16) with Ω set to Q_σ can be used to evaluate the fractional Laplacian.

This shows that the evaluation of the strong form of the fractional Laplacian on all the quadrature nodes $\{\vec{x}_j\}_j$ can be seen as an application of the cluster method with respect to the sets of functions $\{\delta(\cdot - \vec{x}_j)\}_j$ and $\{\phi_i\}_i$.

We illustrate the above results by assembling both the full matrix as well as its sparse approximation on the unit disk for fractional orders $s = 0.25$ and $s = 0.75$. The memory usage of the matrices are compared in Figure 11.5. For low number of degrees of freedom, none of the cluster pairs are admissible, so the full matrix and its approximation have the same size. Starting with roughly 2000 degrees of freedom, the memory footprint of the sparse approximate starts to follow the $n \log^4 n$ curve and therefore outperforms the full assembly. The same behavior can be observed for the assembly times, as seen in Figure 11.6.

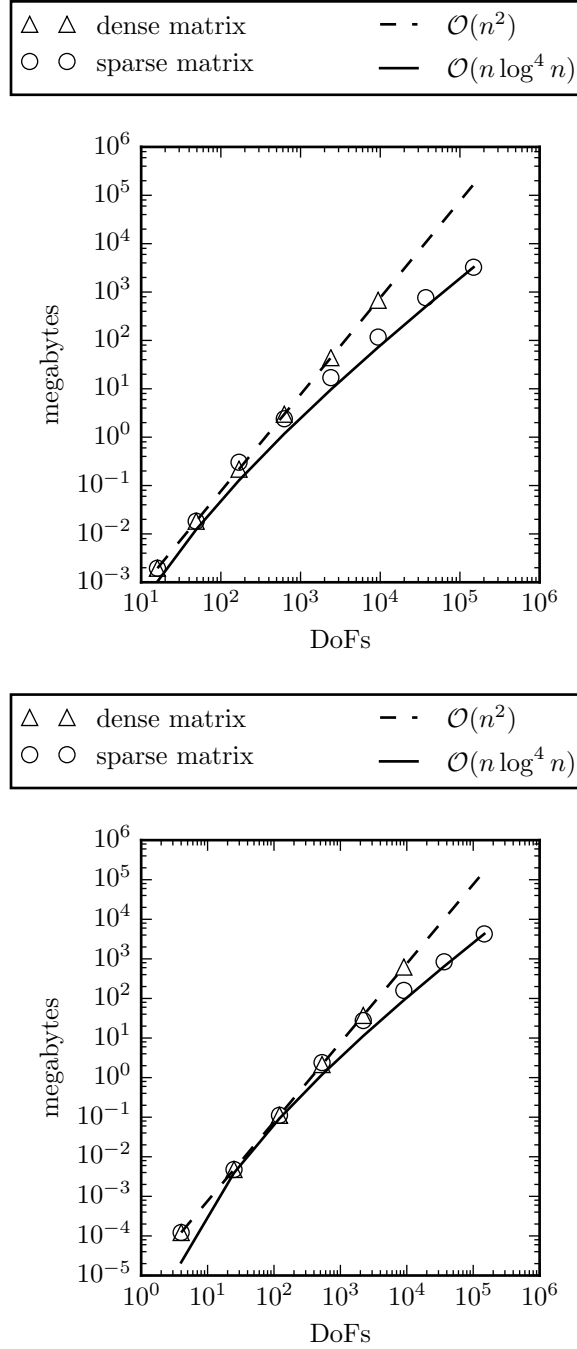


Figure 11.5: Memory usage of the dense matrix and its sparse approximation. $s = 0.25$ (top), $s = 0.75$ (bottom). While the dense matrix uses n^2 floating-point numbers, the sparse approximation can be seen to require only $\mathcal{O}(n \log^4 n)$ memory. At roughly 2000 unknowns, the memory footprint of the sparse approximation separates from the $\mathcal{O}(n^2)$ curve.

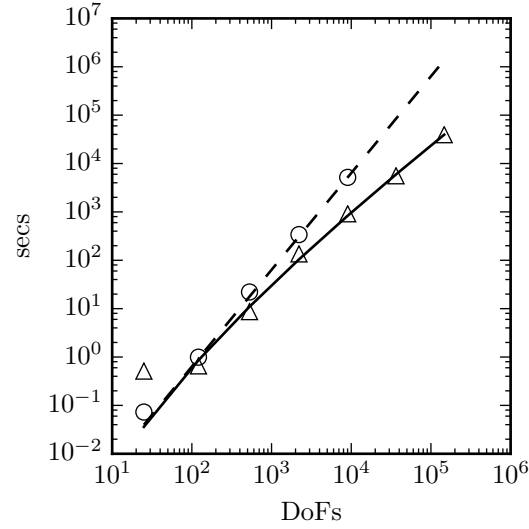
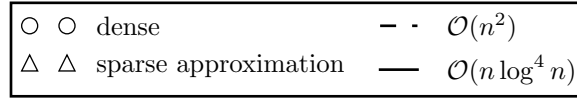
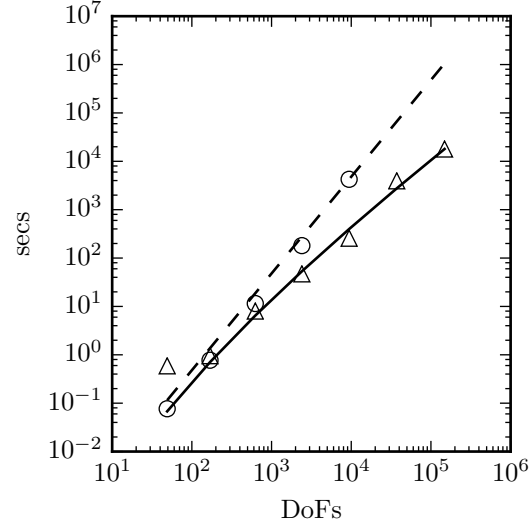
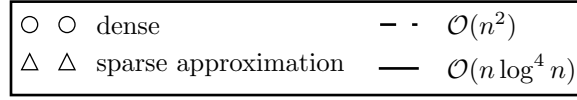


Figure 11.6: Assembly time of the dense matrix and its sparse approximation. $s = 0.25$ (top), $s = 0.75$ (bottom). The time to assemble the full matrix grows quadratically in the number of unknowns, whereas the sparse approximation starts to follow the $n \log^4 n$ curve at about 2000 degrees of freedom.

Chapter 12

Numerical Examples

12.1 Fractional Poisson Equation

For $s \in (0, 1)$, we consider the fractional Poisson problem

$$\begin{aligned} (-\Delta)^s u &= f \text{ in } \Omega = B(0, 1), \\ u &= 0 \text{ in } \Omega^c \end{aligned}$$

in $d \in \{1, 2\}$ dimensions. The reason for choosing a unit disk domain is that analytic solutions to the problem are known. In $d = 1$ dimensions, the right-hand sides

$$f_{n,0}^{1D} = 2^{2s} \Gamma(1+s)^2 \binom{s+n-1/2}{s} \binom{s+n}{s} P_n^{(s,-1/2)}(2x^2-1), \quad n \geq 0$$

have corresponding solutions

$$u_{n,0}^{1D} = P_n^{(s,-1/2)}(2x^2-1) (1-x^2)_+^s$$

and the right-hand sides

$$f_{n,1}^{1D} = 2^{2s} \Gamma(1+s)^2 \binom{s+n+1/2}{s} \binom{s+n}{s} x P_n^{(s,1/2)}(2x^2-1), \quad n \geq 0$$

have solutions

$$u_{n,1}^{1D} = x P_n^{(s,1/2)}(2x^2-1) (1-x^2)_+^s$$

as shown in [45]. Here, $\binom{x}{y} = \frac{\Gamma(x+1)}{\Gamma(y+1)\Gamma(x-y+1)}$ are generalized binomial coefficients, $P_n^{(\alpha,\beta)}$ are the Jacobi polynomials, and $x_+ = \max\{0, x\}$. In $d = 2$ dimensions, the right-hand sides

$$f_{n,\ell}^{2D} = 2^{2s} \Gamma(1+s)^2 \binom{s+n+\ell}{s} \binom{s+n}{s} r^\ell \cos(\ell\theta) P_n^{(s,\ell)}(2r^2-1), \quad \ell, n \geq 0$$

have known analytical solutions given by

$$u_{n,\ell}^{2D} = r^\ell \cos(\ell\theta) P_n^{(s,\ell)}(2r^2-1) (1-r^2)_+^s.$$

For more details on how these analytical solutions are determined, see [45]. We note that both the one-dimensional and the two-dimensional solutions contain a term $(1-x^2)_+^s$ (respectively $(1-r^2)_+^s$) which behaves like $\delta(\vec{x})^s$, where δ is the distance to the boundary.

12.1.1 One-Dimensional Examples

We first solve the fractional Poisson problem with constant right-hand side $f = f_{0,0}^{1D}$ and solution $u = u_{0,0}^{1D}$ for $s = 0.25$ and $s = 0.75$. According to the regularity results in Theorem 34, u is in $H^{s+1/2-\varepsilon}(\Omega)$, so that we expect essentially $h^{1/2} \sim n^{-1/2}$ convergence if a globally quasi-uniform mesh is used. Using graded meshes, the expected rate in $d = 1$ dimensions is n^{s-2} , as shown in eq. (10.11).

We adaptively refine the mesh based on the BR error indicators and a maximum marking strategy with threshold 0.8. The cluster method is used both for the stiffness matrix and for the evaluation of the strong form of the fractional Laplacian. The error plots in Figure 12.1 clearly display n^{s-2} convergence in $\tilde{H}^s$ -norm, both for $s = 0.25$ and for $s = 0.75$, as compared to the $n^{-1/2}$ convergence obtained by uniform mesh refinement. This means that the optimal rate of convergence for $d = 1$ is achieved. The measured L^2 -convergence is of order $n^{-1/2-s}$ for uniform refinement and n^{-2} for adaptive refinement for both values of s .

In Figure 12.2, we plot the number of degrees of freedom n against timing results for the assembly of the stiffness matrix, the solution of the linear system,

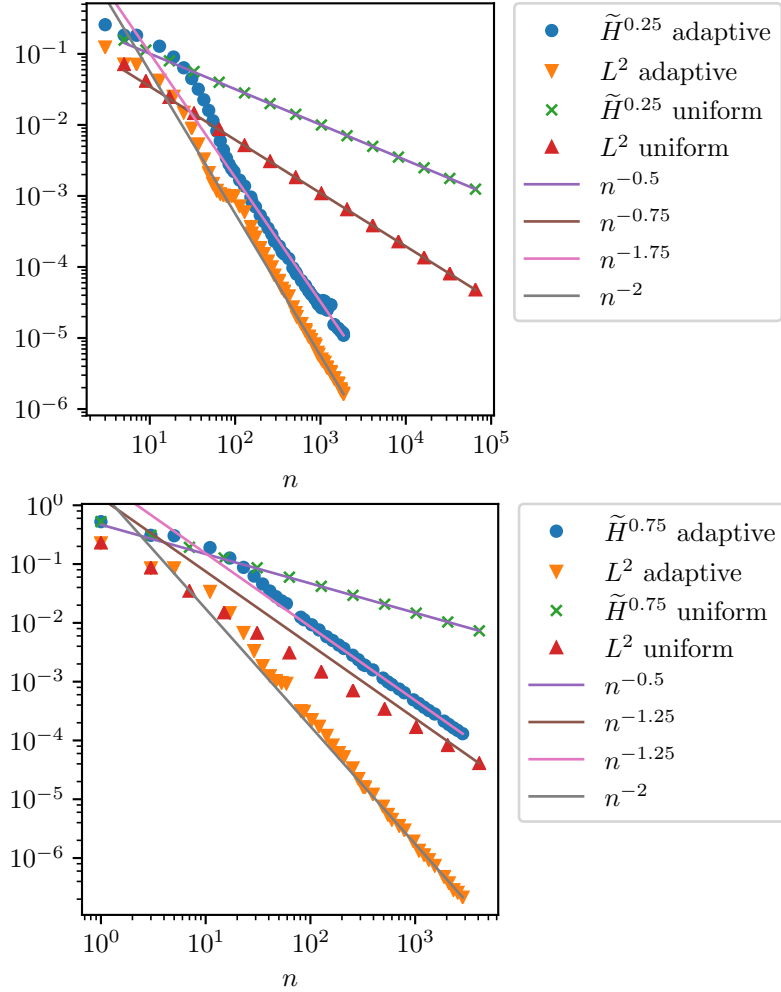


Figure 12.1: $\tilde{H}^s$ - and L^s -error for the one-dimensional problem with constant right-hand side $f = f_{0,0}^{1D}$ and for uniform and adaptive refinement. $s = 0.25$ at the top, $s = 0.75$ at the bottom. It can be seen that while uniform refinement results in convergence in $\tilde{H}^s$ -norm with rate $n^{-1/2}$, adaptive refinement achieves the optimal rate of n^{s-2} . The convergence in L^2 -norm is of order $n^{-1/2-s}$ for uniform refinement, and of order n^{-2} for adaptive refinement.

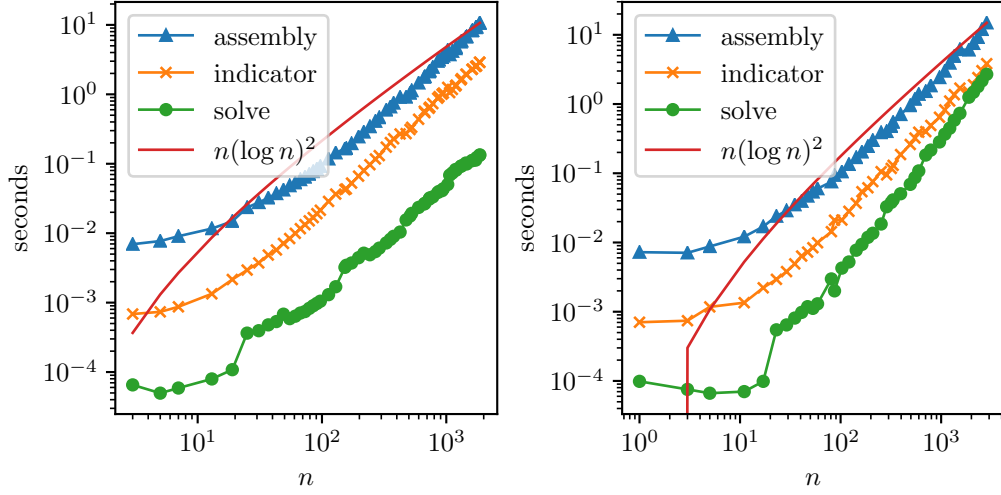


Figure 12.2: Timings for assembly of the stiffness matrix, solution of linear system and computation of the error indicators for the one-dimensional problem with constant right-hand side $f = f_{0,0}^{1D}$. $s = 0.25$ on the left, $s = 0.75$ on the right. All steps can be seen to scale quasi-linearly with the number of unknowns n .

and the computation of the error indicators (10.13). It can be observed that all three scale with $n (\log n)^2$.

Next, we consider the solution of the problem with right-hand side $f(x) = \text{sign } x \in H^{1/2-\varepsilon}(\Omega)$. Based on decomposition of f with respect to $f_{n,\ell}^{1D}$, the analytic solution u is found to be

$$u = 2^{-2s} \sum_{n \geq 0} (-1)^n \frac{2n + s + 3/2}{n + 1/2} \frac{\binom{n+s+1/2}{-1/2}}{\binom{n+s+1/2}{s} \binom{n+s}{s}} x P_n(2x^2 - 1) (1 - x^2)_+^s$$

so that the discretization error can be calculated as

$$\|u - u_h\|_{\tilde{H}^s(\Omega)}^2 = a(u - u_h, u - u_h) = (f, u) - (f, u_h), \quad (12.1)$$

with

$$(f, u) = \frac{2^{1-2s}}{(2s+1)\Gamma(1+s)^2}$$

Solutions for $s = i/10$, $i = 1, \dots, 9$ are plotted in Figure 12.3. For smaller values of s , the transition at the boundary and at $x = 0$ is more pronounced, whereas for

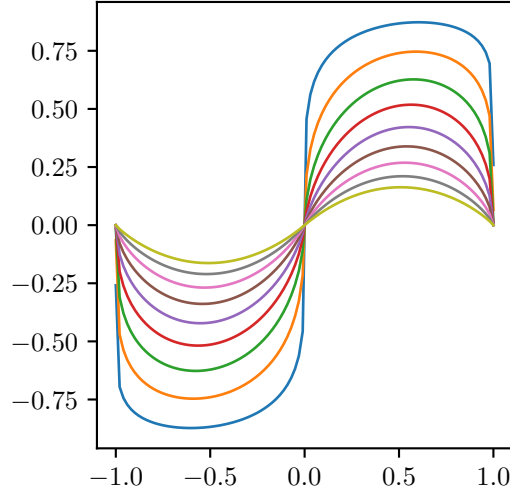


Figure 12.3: Solutions to the one-dimensional fractional Poisson problem with right-hand side $f(x) = \text{sign } x$ for $s = i/10$, $i = 1, \dots, 9$. For smaller values of s , the transition at the boundary and at $x = 0$ is more pronounced, whereas for larger values of s the solutions resemble the solution of the standard integer-order Poisson problem.

larger values of s the solutions resemble the solution of the standard integer-order Poisson problem.

The error plots in Figure 12.4 show that n^{s-2} -convergence in $\tilde{H}^s$ -norm is also achieved for the discontinuous right-hand side $f(x) = \text{sign } x$.

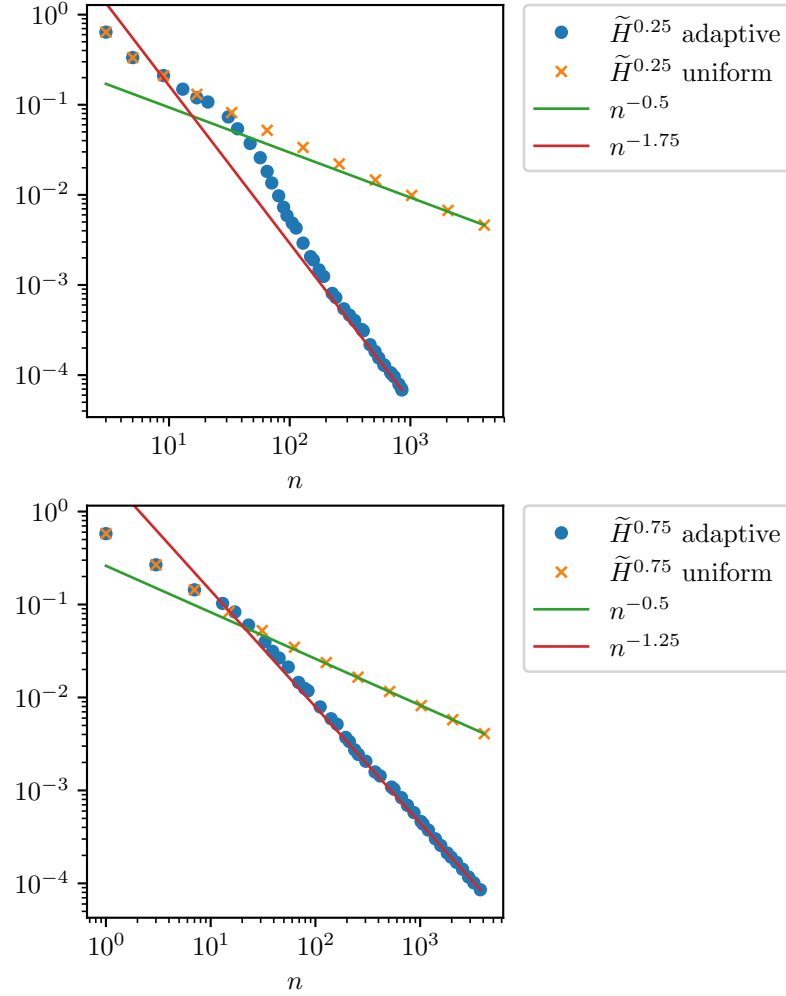


Figure 12.4: $\tilde{H}^s$ -error for the one-dimensional problem with discontinuous right-hand side $f(x) = \text{sign } x$ and for uniform and adaptive refinement. $s = 0.25$ at the top, $s = 0.75$ at the bottom. It can be seen that while uniform refinement results in convergence in $\tilde{H}^s$ -norm with rate $n^{-1/2}$, adaptive refinement achieves the optimal rate of n^{s-2} .

12.1.2 Two-Dimensional Examples

We consider the two-dimensional fractional Poisson problem on a disk with constant right-hand side $f = f_{0,0}^{2D}$. The solutions u for $s = 0.25$ and $s = 0.75$ are shown in Figure 10.1.

Using a globally quasi-uniform mesh, we would obtain $h^{1/2} \sim n^{-1/4}$ convergence. Using a graded mesh, $n^{-1/2}$ convergence was predicted, owing to the fact that the mesh is required to be locally quasi-uniform. Again, we adaptively refine the mesh, based on the error indicators (10.13). The obtained meshes (Figure 12.5) are highly refined in regions close to the boundary, where the solution is singular.

We plot the $\tilde{H}^s(\Omega)$ -error in Figure 12.6. It can be observed that $n^{-1/2}$ convergence is obtained, thus matching the optimal rate obtained using a graded mesh for a constant right-hand side in (10.12). As in the one-dimensional case, the rate in L^2 -norm is s orders higher than the rate in $\tilde{H}^s$ -norm. By considering the timing results shown in Figure 12.7, we see that all phases of the adaptive refinement loop scale with $n(\log n)^4$. The computation of the error indicators is cheaper than the assembly of the system matrix, and the solution of the linear system is significantly cheaper than both. Relative to the matrix assembly, the computation of the error indicators is more expensive than for the one dimensional problem, since it involves numerical quadrature over element edges instead of pointwise evaluations.

We now consider the solution of the problem with right-hand side $f(r, \theta) = 1_{|\theta| < \pi/2} \in H^{1/2-\varepsilon}(\Omega)$. Based on decomposition of f with respect to $f_{n,\ell}^{2D}$, the analytic solution u is found to be

$$u = 2^{-2s} \frac{(1-r^2)_+^s}{\Gamma(1+s)^2} \left\{ \frac{1}{2} + \sum_{n \geq 0, \ell \geq 1 \text{ and odd}} \frac{(-1)^{(\ell+1)/2+n} (2n+s+l+1) \cos(\ell\theta) r^\ell P_n^{(s,\ell)}(2r^2-1)}{\pi(n+l/2)(s+1) \binom{n+s+l/2+1}{n+l/2} \binom{s+n}{n}} \right\},$$

so that the discretization error can be calculated as

$$\|u - u_h\|_{\tilde{H}^s(\Omega)}^2 = a(u - u_h, u - u_h) = (f, u) - (f, u_h),$$

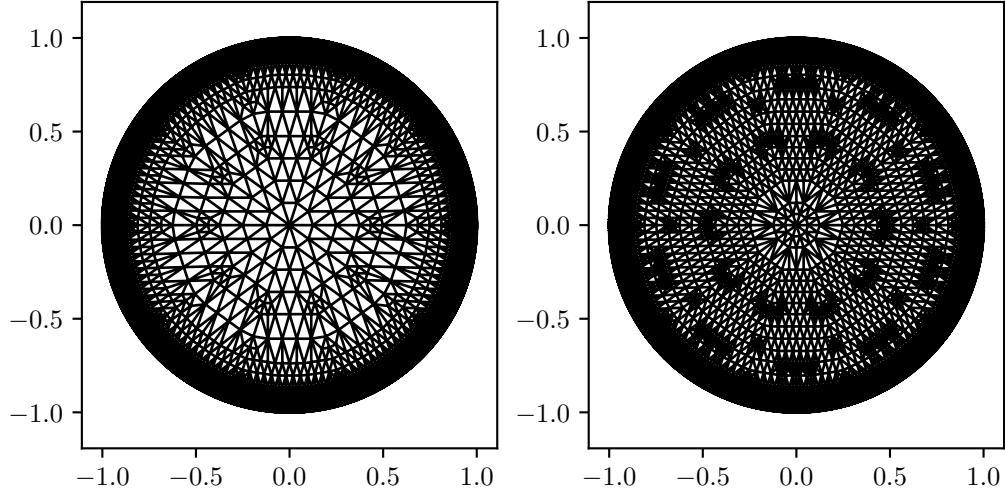


Figure 12.5: Adaptively refines meshes for the two-dimensional problem with constant right-hand side. $s = 0.25$ on the left, $s = 0.75$ on the right.

with

$$(f, u) = 2^{-2s} \left\{ \frac{\pi}{4(s+1)\Gamma(s+1)^2} - \frac{1}{\pi} G_{4,4}^{3,2} \left(\begin{matrix} 1, & 1+s/2, & 5/2+s, & 5/2+s \\ 2, & 1/2, & 1/2, & 2+s/2 \end{matrix} \middle| -1 \right) \right\},$$

where G is the Meijer G-function [99].

The solutions u for $s = 0.25$ and $s = 0.75$ are shown in Figure 12.8. The obtained meshes (Figure 12.9) are highly refined in regions close to the boundary and along the discontinuity $\theta = \pm\pi/2$. We plot the $\tilde{H}^s$ -error in Figure 12.10. In the case $s = 0.25$, we recover the optimal rate of convergence of $n^{-1/2}$, whereas for $s = 0.75$, we only find $n^{-1/4}$, which coincides with the rate that would be obtained by uniform refinement.

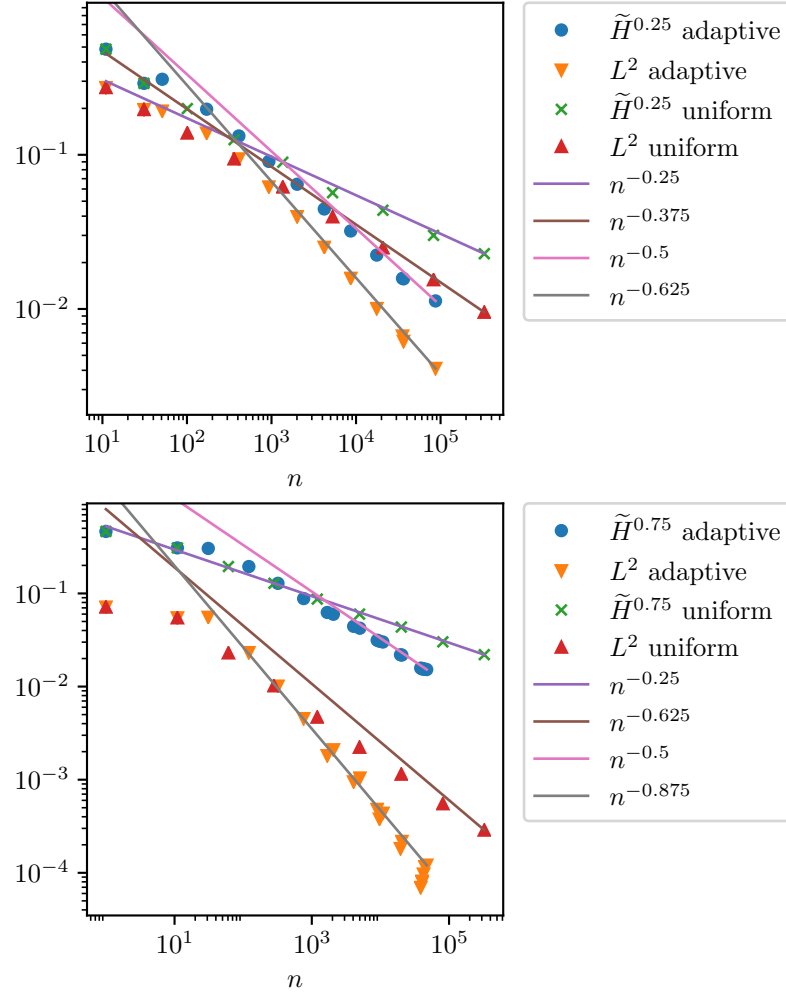


Figure 12.6: $\tilde{H}^s$ - and L^s -error for the two-dimensional problem with constant right-hand side $f = f_{0,0}^{2D}$ and for uniform and adaptive refinement. $s = 0.25$ at the top, $s = 0.75$ at the bottom. It can be seen that while uniform refinement results in convergence in $\tilde{H}^s$ -norm with rate $n^{-1/4}$, adaptive refinement achieves the optimal rate of $n^{-1/2}$. The convergence in L^2 -norm is of order $n^{-1/4-s/2}$ for uniform refinement, and of order $n^{-1/2-s/2}$ for adaptive refinement.

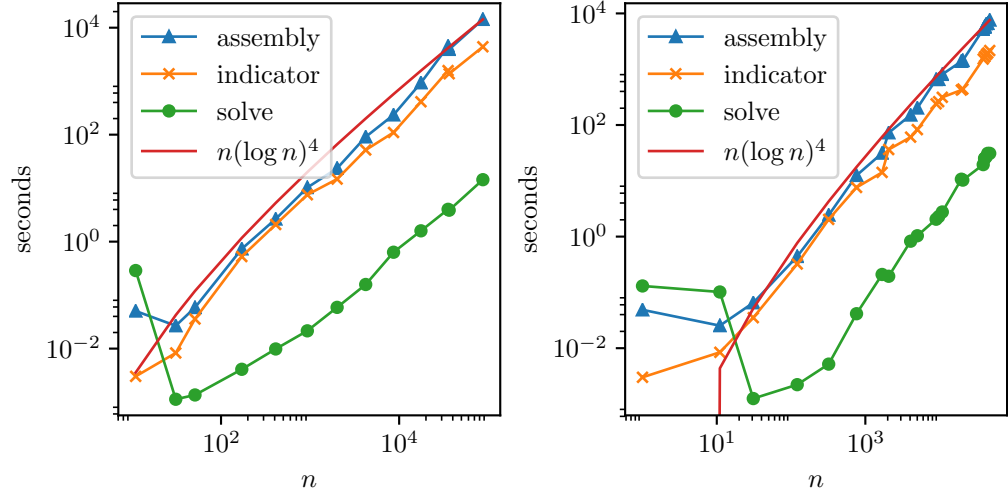


Figure 12.7: Timings for assembly of the stiffness matrix, solution of linear system and computation of the error indicators for the two-dimensional problem with constant right-hand side $f = f_{0,0}^{2D}$. $s = 0.25$ on the left, $s = 0.75$ on the right. All steps can be seen to scale quasi-linearly with the number of unknowns n .

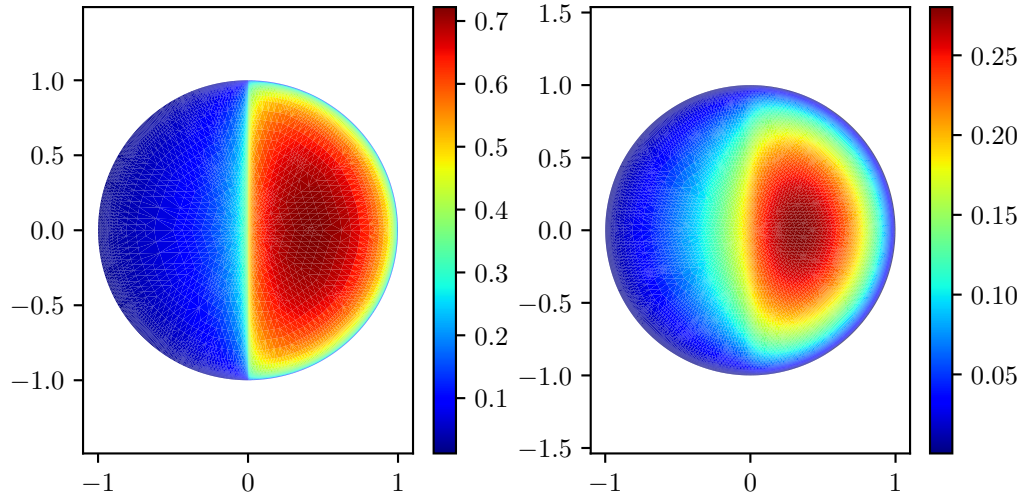


Figure 12.8: Solutions u corresponding to the discontinuous right-hand side $f(r, \theta) = 1_{|\theta| < \pi/2}$ for $s = 0.25$ and for $s = 0.75$.

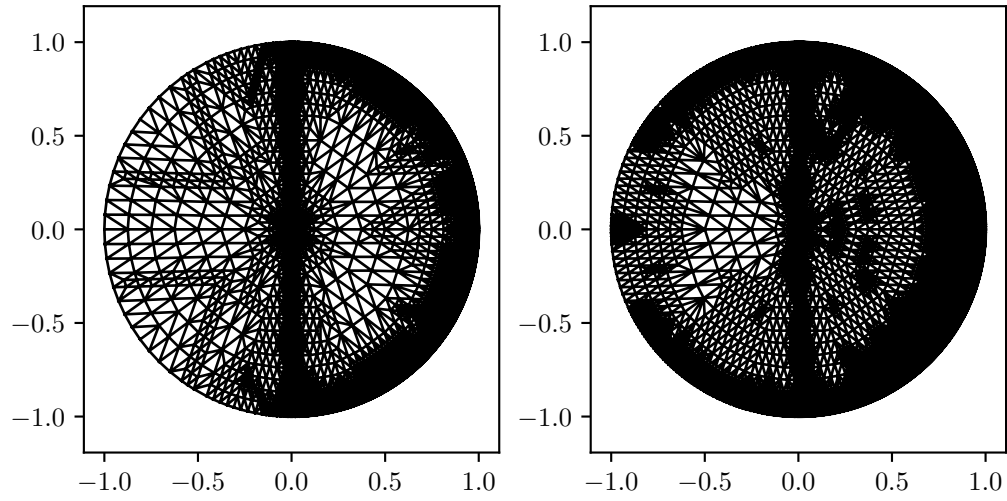


Figure 12.9: Adaptively refined meshes for the discontinuous right-hand side. $s = 0.25$ on the left, $s = 0.75$ on the right.

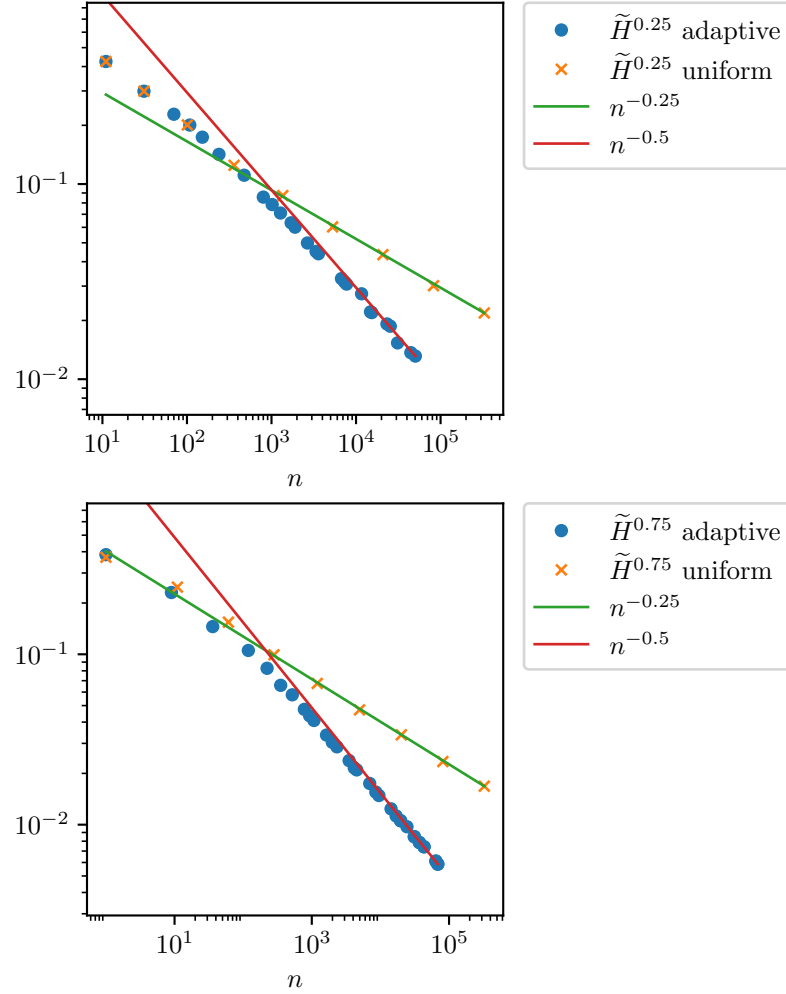


Figure 12.10: $\tilde{H}^s$ -error for the two-dimensional problem with discontinuous right-hand side $f(r, \theta) = 1_{|\theta| < \pi/2}$ and for uniform and adaptive refinement. $s = 0.25$ at the top, $s = 0.75$ at the bottom. For $s = 0.25$, uniform refinement results in convergence in $\tilde{H}^s$ -norm with rate $n^{-1/2}$, adaptive refinement achieves the optimal rate of n^{s-2} . For $s = 0.75$, both uniform and adaptive refinement results in convergence in $\tilde{H}^s$ -norm with rate $n^{-1/2}$.

12.1.3 L-shaped Domain

We propose to solve the fractional Poisson problem on the L-shaped domain $\Omega = [0, 2]^2 \setminus [1, 2]^2$ with constant right-hand side. Since no analytic solution is available, we solve the problem on a highly refined mesh and use this reference solution $u_{\underline{h}}$ to approximate the error. In Figure 12.13, we show $\tilde{H}^s$ - and L^2 error, as well as the a posteriori error estimate $\hat{\eta} = (\sum_i \hat{\eta}_i^2)^{1/2}$, which is based on the local a posteriori error estimators $\hat{\eta}_i$ given in (10.13). It can be observed that $\|u_h - u_{\underline{h}}\|_{\tilde{H}^s}$ and the a posteriori error estimate η converge with optimal rate $n^{-1/2}$, and that $\|u_h - u_{\underline{h}}\|_{L^2}$ converges with rate $n^{-1/2-s/2}$. The apparent speed-up of convergence in $\tilde{H}^s$ - and L^2 -norm for larger number of unknowns is due to the fact that the reference solution $u_{\underline{h}}$ is used in their computation instead of the true solution u . The a posteriori error estimate is obviously not suffer this deficiency, because it does not involve the reference solution $u_{\underline{h}}$. This example demonstrates that more complicated domains can easily be considered in the presented framework, and that convergence of optimal order can equally well be achieved.

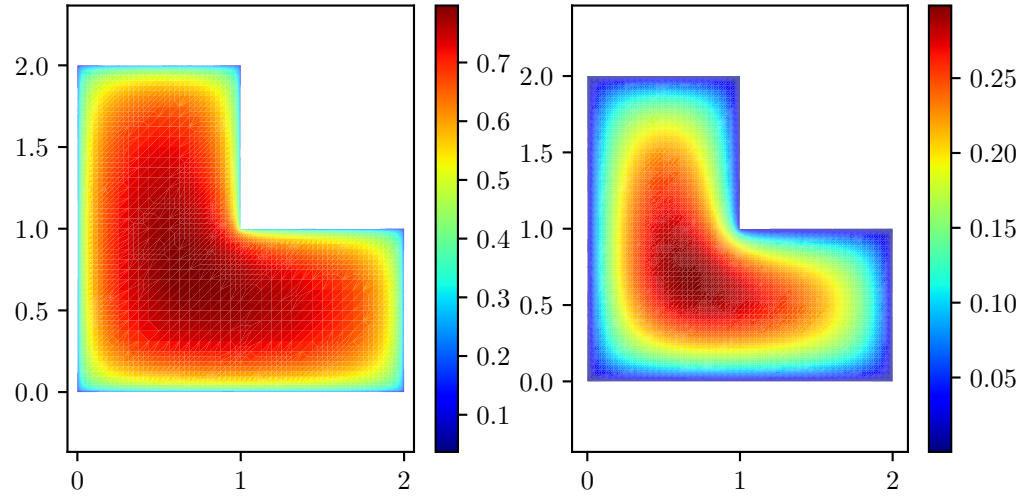


Figure 12.11: Solutions u on the L-shaped domain corresponding to a constant right-hand side for $s = 0.25$ and for $s = 0.75$.

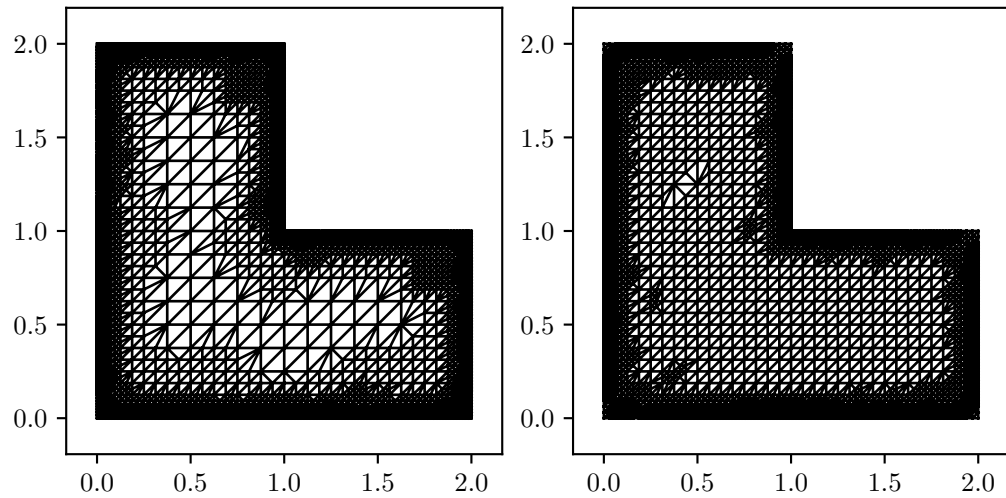


Figure 12.12: Adaptively refined meshes for the L-shaped domain. $s = 0.25$ on the left, $s = 0.75$ on the right.

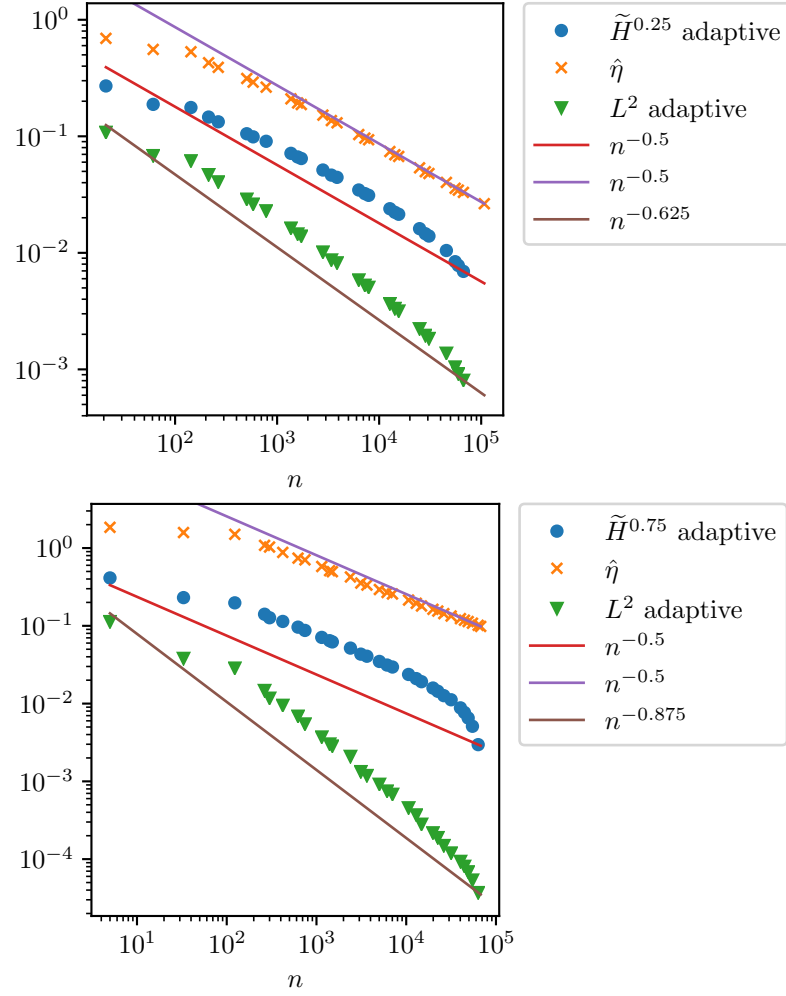


Figure 12.13: $\tilde{H}^s$ -, L^2 -error and estimated error $\hat{\eta}$ for the L-shaped domain. $s = 0.25$ at the top, $s = 0.75$ at the bottom. Adaptive refinement results in convergence in $\tilde{H}^s$ -norm with optimal rate of $n^{-1/2}$, and the same holds for the a posteriori error $\hat{\eta}$. In L^2 -norm, the rate $n^{-1/2-s/2}$ is observed.

12.1.4 Exploiting Symmetry

The example of the previous sections displayed symmetries and periodicity in the shape of the domain and the right-hand side and therefore in the solution. In particular, the solution shown in Figure 12.8 suggests that one should be able to exploit symmetry to halve the number of degrees of freedom required.

Let $\mathbf{B}$ be an orthogonal transformation such that for some K , $\mathbf{B}^K = \text{Id}$, $\mathbf{B}\Omega = \Omega$, and let $\omega \subset \tilde{\omega} \subset \mathbb{R}^d$ be representative domains such that $\mathbb{R}^d = \bigcup_{i=0}^{K-1} \mathbf{B}^i \tilde{\omega}$ and $\Omega = \bigcup_{i=0}^{K-1} \mathbf{B}^i \omega$. Moreover, assume that $u(\mathbf{B}^i \vec{x}) = u(\vec{x})$ and $f(\mathbf{B}^i \vec{x}) = f(\vec{x})$. Examples of such transformations for $\Omega = B(0, 1)$ include the reflection $\mathbf{B} = -\text{Id}$ with $\tilde{\omega}$ and ω being the half-space and the half disc, or the rotation by angle $\theta = \frac{2\pi}{K}$ and the domains $\tilde{\omega} = \{(r \cos \phi, r \sin \phi) \in \mathbb{R}^2 \mid r \geq 0, \phi \in [0, 2\pi/K]\}$ and $\omega = \tilde{\omega} \cap \{r \leq 1\}$.

The variational formulation of the fractional Poisson problem eq. (10.3) then reduces to

$$\begin{aligned} a(u, v) &= \frac{C(d, s)}{2} \int_{\omega} \int_{\omega} \sum_{i,j=0}^{K-1} \frac{(u(\mathbf{B}^i \vec{x}) - u(\mathbf{B}^j \vec{y})) (v(\mathbf{B}^i \vec{x}) - v(\mathbf{B}^j \vec{y}))}{|\mathbf{B}^i \vec{x} - \mathbf{B}^j \vec{y}|^{d+2s}} \\ &\quad + C(d, s) \int_{\omega} \int_{\tilde{\omega} \setminus \omega} \sum_{i,j=0}^{K-1} \frac{u(\mathbf{B}^i \vec{x}) v(\mathbf{B}^j \vec{x})}{|\mathbf{B}^i \vec{x} - \mathbf{B}^j \vec{y}|^{d+2s}} \\ &= K \frac{C(d, s)}{2} \int_{\omega} \int_{\omega} (u(\vec{x}) - u(\vec{y})) (v(\vec{x}) - v(\vec{y})) k_{\text{eff}}(\vec{x}, \vec{y}) \\ &\quad + KC(d, s) \int_{\omega} \int_{\tilde{\omega} \setminus \omega} u(\vec{x}) v(\vec{x}) k_{\text{eff}}(\vec{x}, \vec{y}) \end{aligned}$$

and

$$\begin{aligned} \langle f, v \rangle &= \sum_{i=0}^{K-1} \int_{\omega} f(\mathbf{B}^i \vec{x}) v(\mathbf{B}^i \vec{x}) \\ &= K \int_{\omega} f v \end{aligned}$$

with the effective kernel

$$k_{\text{eff}}(\vec{x}, \vec{y}) = \sum_{i=0}^{K-1} |\vec{x} - \mathbf{B}^i \vec{y}|^{-d-2s}.$$

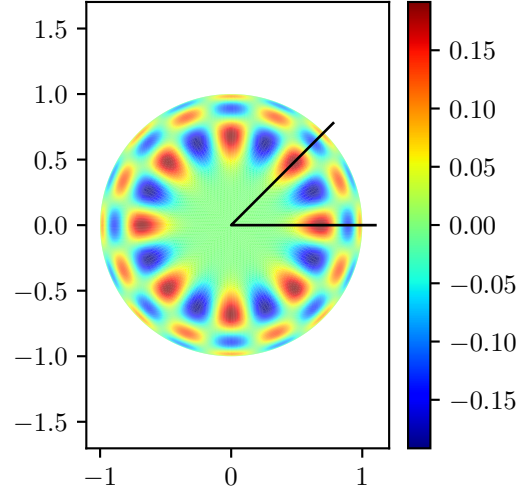


Figure 12.14: Periodic solution $u = u_{2,8}^{2D}$. Only the circular sector on the right is discretized, with periodic images taken into account through the effective kernel k_{eff} .

that accounts for the interaction with mirror images. Effectively, the computational effort for the assembly and solution of systems involving the dense stiffness matrix is divided by a factor of K , since only the sub-domain ω needs to be discretized. The complexity of the cluster method is expected to decrease from $\mathcal{O}\left(n(\log n)^{2d}\right)$ to $\mathcal{O}\left(n/K(\log n)^{2d}\right)$ since the assembly of the near-field contributions constitutes the bulk of the work.

To illustrate the approach, we solve the fractional Poisson problem with right-hand side $f_{2,8}$ on the unit disk. By leveraging the periodicity in the angular direction (and neglecting the mirror symmetry), we reduce to a problem on the circular sector $\tilde{\omega} = \{(r \cos \phi, r \sin \phi) \in \mathbb{R}^2 \mid r \geq 0, \phi \in [0, \pi/4]\}$ and $\omega = \tilde{\omega} \cap \{r \leq 1\}$. For simplicity, we employ a uniform mesh and assemble the dense system matrix. The solution is shown in Figure 12.14.

12.2 Fractional Heat Equation

The fractional heat equation is given by

$$\begin{aligned} u_t + (-\Delta)^s u &= f && \text{in } \Omega, \\ u &= 0 && \text{in } \Omega^c. \end{aligned}$$

We propose to approximate the problem using an implicit method in time. The simplest such scheme is the backward Euler method

$$(\mathbf{M} + \Delta t \mathbf{A}^s) \vec{u}^{k+1} = \mathbf{M} \vec{u}^k + \vec{f}^{k+1},$$

where $u(\cdot, k\Delta t) \approx \sum_i u_i^k \phi_i$ and $f_i^k = (f(\cdot, k\Delta t), \phi_i)$.

More generally, let us assume that a scheme of order α is used in time, and that a globally quasi-uniform mesh with mesh size h is used in space. In order to obtain optimal convergence in L^2 -norm, in view of Theorem 37, we shall choose $\Delta t^\alpha \sim h^{1/2+\min(1/2,s)}$, i.e.

$$\Delta t_{L^2} \sim h^{\min(2,1+2s)/(2\alpha)}.$$

On the other hand, if optimal $\tilde{H}^s(\Omega)$ -convergence is desired, we need $\Delta t_{\tilde{H}^s(\Omega)} \sim h^{1/(2\alpha)}$, see Lemma 36. Consequently, if an order α scheme is used for time stepping, with optimal time step Δt_{L^2} or $\Delta t_{\tilde{H}^s(\Omega)}$, we find by Lemma 42 that the condition numbers of the iteration matrix satisfy

$$\begin{aligned} \kappa(\mathbf{M} + \Delta t_{L^2} \mathbf{A}^s) &\leq C (1 + h^{\min(2,1+2s)/(2\alpha)-2s}), \\ \kappa(\mathbf{M} + \Delta t_{\tilde{H}^s(\Omega)} \mathbf{A}^s) &\leq C (1 + h^{1/(2\alpha)-2s}). \end{aligned}$$

In particular, this shows that the condition number will not grow as the mesh size decreases if $s \leq 1/(4\alpha - 2)$ in the L^2 case, or if $s \leq 1/(4\alpha)$ in the $\tilde{H}^s(\Omega)$ case.

We illustrate the consequences of the above result in the case of a second order accurate time stepping scheme ($\alpha = 2$), and for $s = 0.25$ and $s = 0.75$. In the case of $s = 0.25$, $\Delta t_{L^2} \sim h^{3/8}$ and $\kappa(\mathbf{M} + \Delta t_{L^2} \mathbf{A}^s) \sim 1 + h^{-1/8}$. This suggests

Method	$\Delta t = \Delta t_{L^2}$	$\Delta t = \Delta t_{\tilde{H}^s(\Omega)}$
Conjugate Gradient	$n^{1+2s/d-\min(2,1+2s)/(2\alpha d)} (\log n)^{2d}$	$n^{1+2s/d-1/(2\alpha d)} (\log n)^{2d}$
Multigrid	$n (\log n)^{2d}$	$n (\log n)^{2d}$

Table 12.1: Complexity of different solvers for $(\mathbf{M} + \Delta t \mathbf{A}^s) \vec{u} = \vec{b}$ for $\Delta t = \Delta t_{L^2}$ and $\Delta t = \Delta t_{\tilde{H}^s(\Omega)}$ for an α -order time stepping scheme.

that the conjugate gradient method will uniformly in h deliver good results. The convergence of the multigrid method does not depend on the condition number and is essentially independent of h . This is indeed what is observed in the top part of Figure 12.15. In Figure 12.16, the number of iterations is shown. It can be observed that for $s = 0.25$ both the multigrid and the conjugate gradient solver require an essentially constant number of iterations for varying values of Δt .

On the other hand, for $s = 0.75$, $\Delta t_{L^2} \sim h^{1/2}$ and $\kappa(\mathbf{M} + \Delta t_{L^2} \mathbf{A}^s) \sim 1 + h^{-1}$. Therefore, the condition number increases as h goes to zero, and we expect that multigrid asymptotically outperforms the CG solver. This is indeed what is observed in Figures 12.15 and 12.16.

The complexities of the different solvers for different choices of time step size are summarized in Table 12.1.

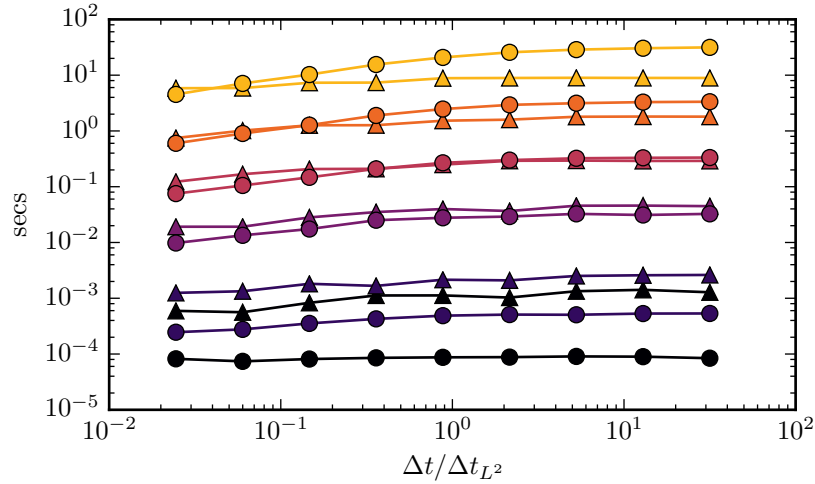
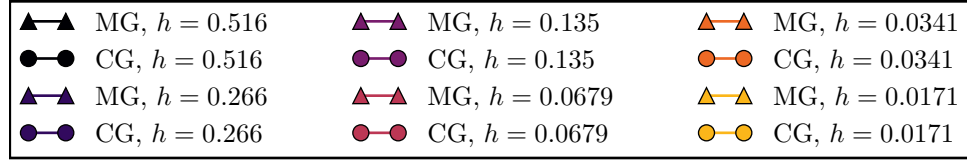
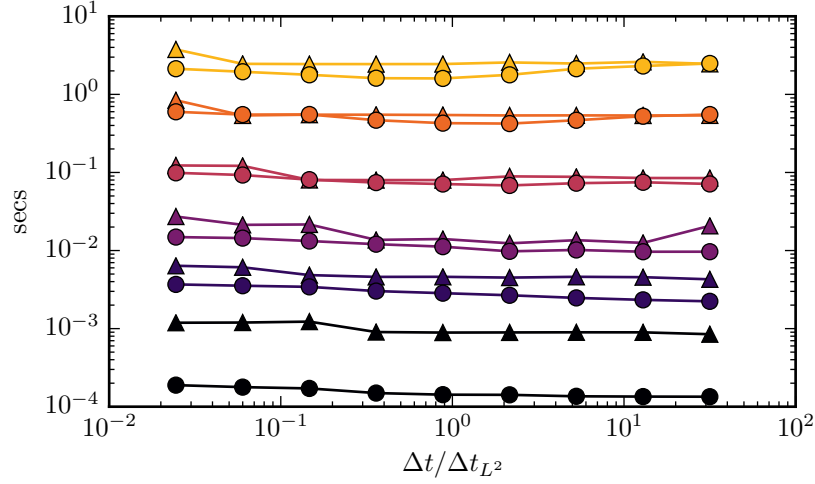
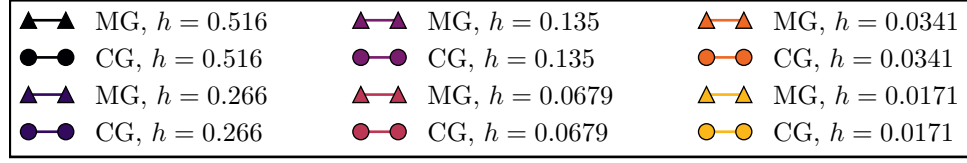


Figure 12.15: Timings in seconds for CG and MG depending on Δt for $s = 0.25$ (top) and $s = 0.75$ (bottom). It can be observed that, for $s = 0.25$, the conjugate gradient method is essentially on par with the multigrid solver. For $s = 0.75$, the multigrid solver asymptotically outperforms the conjugate gradient method, since the condition number $\kappa(M + \Delta t_{L^2} A^s)$ grows as h^{-1} .

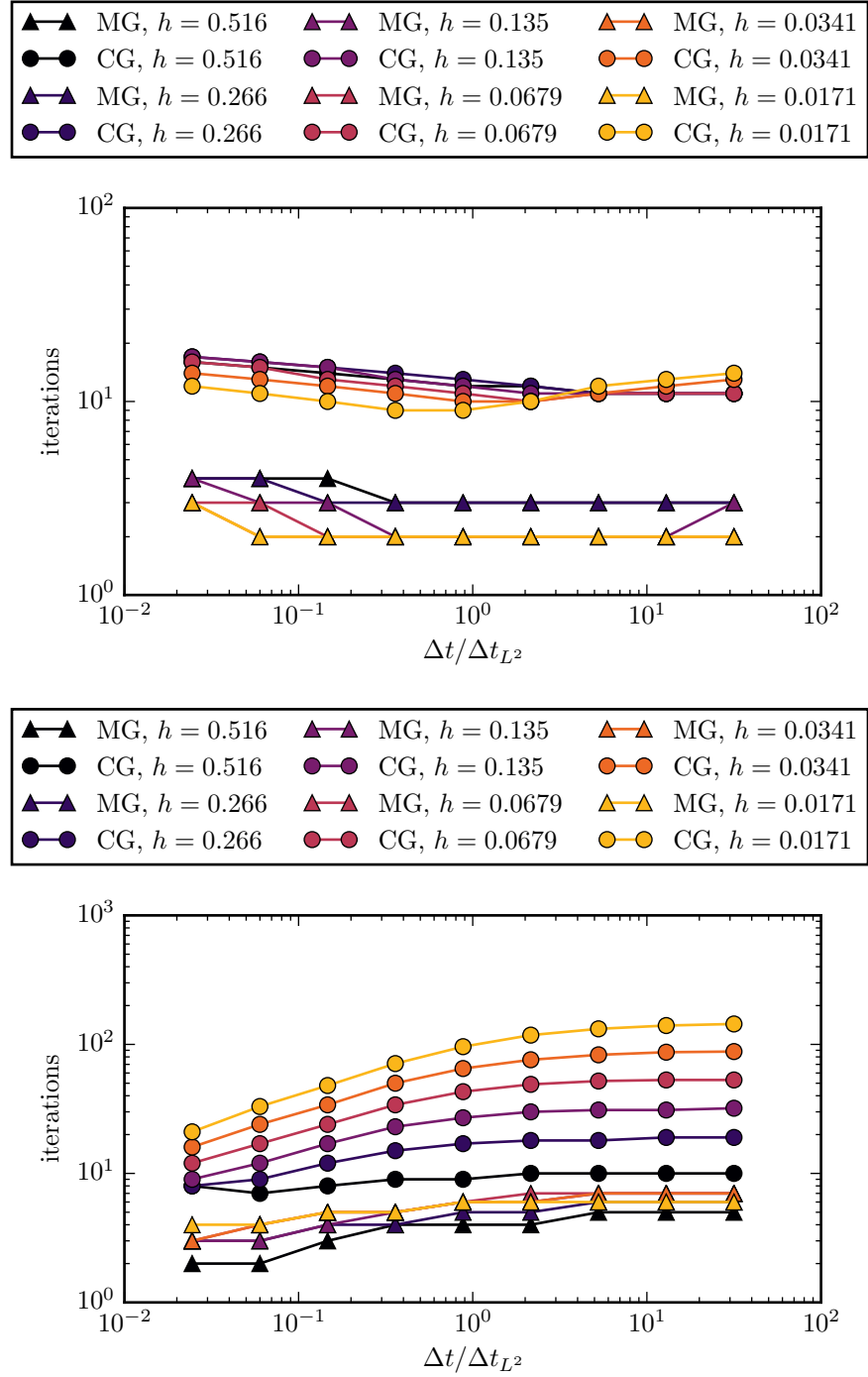


Figure 12.16: Number of iterations for CG and MG depending on Δt for $s = 0.25$ (top) and $s = 0.75$ (bottom). For $s = 0.25$, the number of iterations is essentially independent of Δt . For $s = 0.75$, the number of iterations of the multigrid solver is independent of Δt , but the iterations count for conjugate gradient grows with $h^{-1/2}$.

12.3 Fractional Reaction-Diffusion Systems

In [60], a space-fractional Brusselator model was analyzed and compared to the classical integer-order case. The coupled system of equations is given by

$$\begin{aligned}\frac{\partial X}{\partial t} &= -D_X (-\Delta)^\alpha X + A - (B+1)X + X^2Y, \\ \frac{\partial Y}{\partial t} &= -D_Y (-\Delta)^\beta Y + BX - X^2Y.\end{aligned}$$

Here, D_X and D_Y are diffusion coefficients, A and B are reaction parameters, and α and β determine the type of diffusion. By rewriting the solutions as deviations from the stationary solution $X = A$, $Y = B/A$ and rescaling, one obtains

$$\frac{\partial u}{\partial t} = -(-\Delta)^\alpha u + (B-1)u + Q^2v + \frac{B}{Q}u^2 + 2Quv + u^2v, \quad (12.2)$$

$$\eta^2 \frac{\partial v}{\partial t} = -(-\Delta)^\beta v - Bu - Q^2v - \frac{B}{Q}u^2 - 2Quv - u^2v, \quad (12.3)$$

with $\eta = \sqrt{D_Y/D_X^{\beta/\alpha}}$ and $Q = A\eta$.

In [60] the equations were augmented with periodic boundary conditions and approximated using a pseudospectral method for various different parameter combinations. Here, thanks to the foregoing developments, we have the flexibility to handle more general domains and, in particular, we consider the case where Ω corresponds to a Petri-dish. We solve the above set of equations using a second order accurate IMEX scheme proposed by Koto [72], whose Butcher tableaux are given by Table 12.2. The diffusive parts are treated implicitly and therefore require the solution of several systems all of which are of the type $M + c\Delta t A^s$ with appropriate values of c .

0	0				0				
1	0	1			1	1			
1/2	0	-1/2	1		1/2	0	0		
1	0	-1	1	1	1	0	0	1	
	0	-1	1	1		0	0	1	0

Table 12.2: IMEX scheme by Koto. Implicit scheme on the left, explicit on the right.

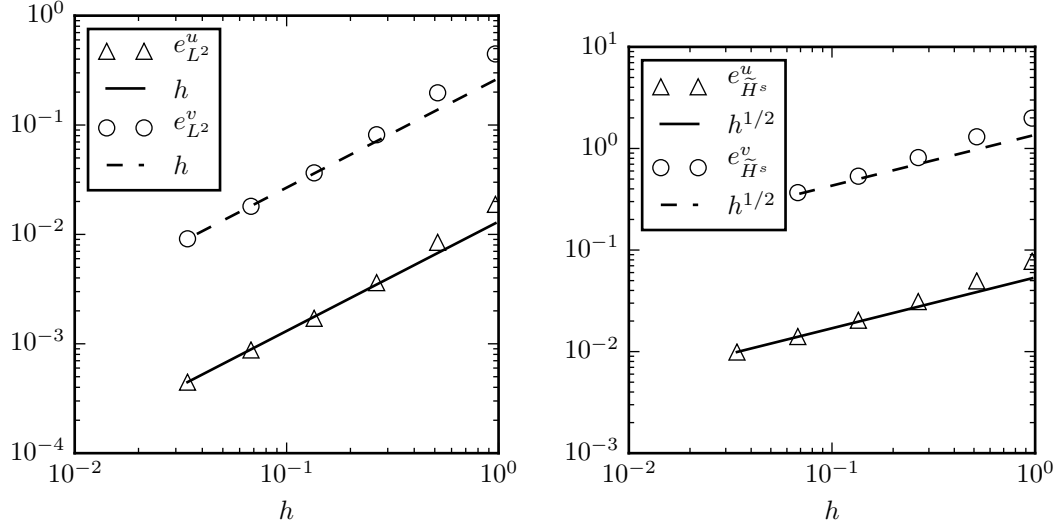


Figure 12.17: Error in L^2 -norm (top) and $\tilde{H}^s(\Omega)$ -norm (bottom) in the Brusselator model. Optimal orders of convergence are achieved. (Compare Theorem 37 and lemma 36.)

In order to verify the correct convergence behavior, we add forcing functions f and g to the system, chosen such that the analytic solution is given by

$$\begin{aligned} u &= \eta \sin(t) u^s(\vec{x}), \\ v &= \eta^{-1} \cos(2t) u^s(\vec{x}), \end{aligned}$$

for suitable initial conditions, where u^s is the solution of the fractional Poisson problem with constant right-hand side. We take $\alpha = \beta = 0.75$, and choose $\Delta t \sim h^{1/2}$, since we already saw that the rate of the spatial approximation in L^2 -norm is of order h . We measure the error as

$$\begin{aligned} e_{L^2}^u &= \max_{0 \leq t_i \leq 10} \|u(t_i, \cdot) - u_h^i\|_{L^2}, & e_{L^2}^v &= \max_{0 \leq t_i \leq 10} \|v(t_i, \cdot) - v_h^i\|_{L^2}, \\ e_{\tilde{H}^s(\Omega)}^u &= \max_{0 \leq t_i \leq 10} \|u(t_i, \cdot) - u_h^i\|_{\tilde{H}^s(\Omega)}, & e_{\tilde{H}^s(\Omega)}^v &= \max_{0 \leq t_i \leq 10} \|v(t_i, \cdot) - v_h^i\|_{\tilde{H}^s(\Omega)}. \end{aligned}$$

From the error plots in Figure 12.17, it can be observed that $e_{L^2} \sim h$ and $e_V \sim h^{1/2}$, as expected.

Having verified the accuracy of the method, we turn to the solution of the system eqs. (12.2) and (12.3) augmented with exterior Neumann conditions as described

in Section 10.2. Golovin, Matkowsky and Volpert [60] observed that for $\eta = 0.2$, $B = 1.22$ and $Q = 0.1$, a single localized perturbation would first form a ring and then break up into spots. The radius of the ring and the number of resulting spots increases as the fractional orders are decreased. In Figure 12.18, simulation results for $\alpha = \beta = 0.625$ and $\alpha = \beta = 0.75$ are shown. We observe that in both cases, an initially circular perturbation develops into a ring. Lower diffusion coefficients do lead to a larger ring, which breaks up later and into more spots. In the last row, we can see that the resulting spots start to replicate and spread out over the whole domain.

Another choice of parameters leads to stripes in the solution. For $\alpha = \beta = 0.75$, $\eta = 0.2$, $B = 6.26$ and $Q = 2.5$, and a random initial condition, stripes without directionality form in the whole domain. This is in alignment with the theoretical considerations of Golovin, Matkowsky and Volpert [60].

12.4 Discussion

In this work, we have presented the numerical approximation of the fractional Laplacian operator for the case of globally quasi-uniform meshes as well as locally refined meshes. In order to overcome slow convergence due to the inherent singularity of the solution close to the boundary, we introduced a posteriori error indicators and adaptive mesh refinement. The non-locality of the problem leads to dense matrices and solution complexity of $\mathcal{O}(n^3)$ using direct solvers and at best $\mathcal{O}(n^2)$ using an optimal iterative method. We showed that by employing a cluster method, assembly, matrix-vector product and computation of the error estimators scale quasi-linear in the number of unknowns. This means that the solution of the fractional Poisson problem (10.1) and the fractional Heat equation (10.2) can be performed in the same complexity and with equal (or even better) rate of convergence as of the standard integer order equivalents. Through several one- and

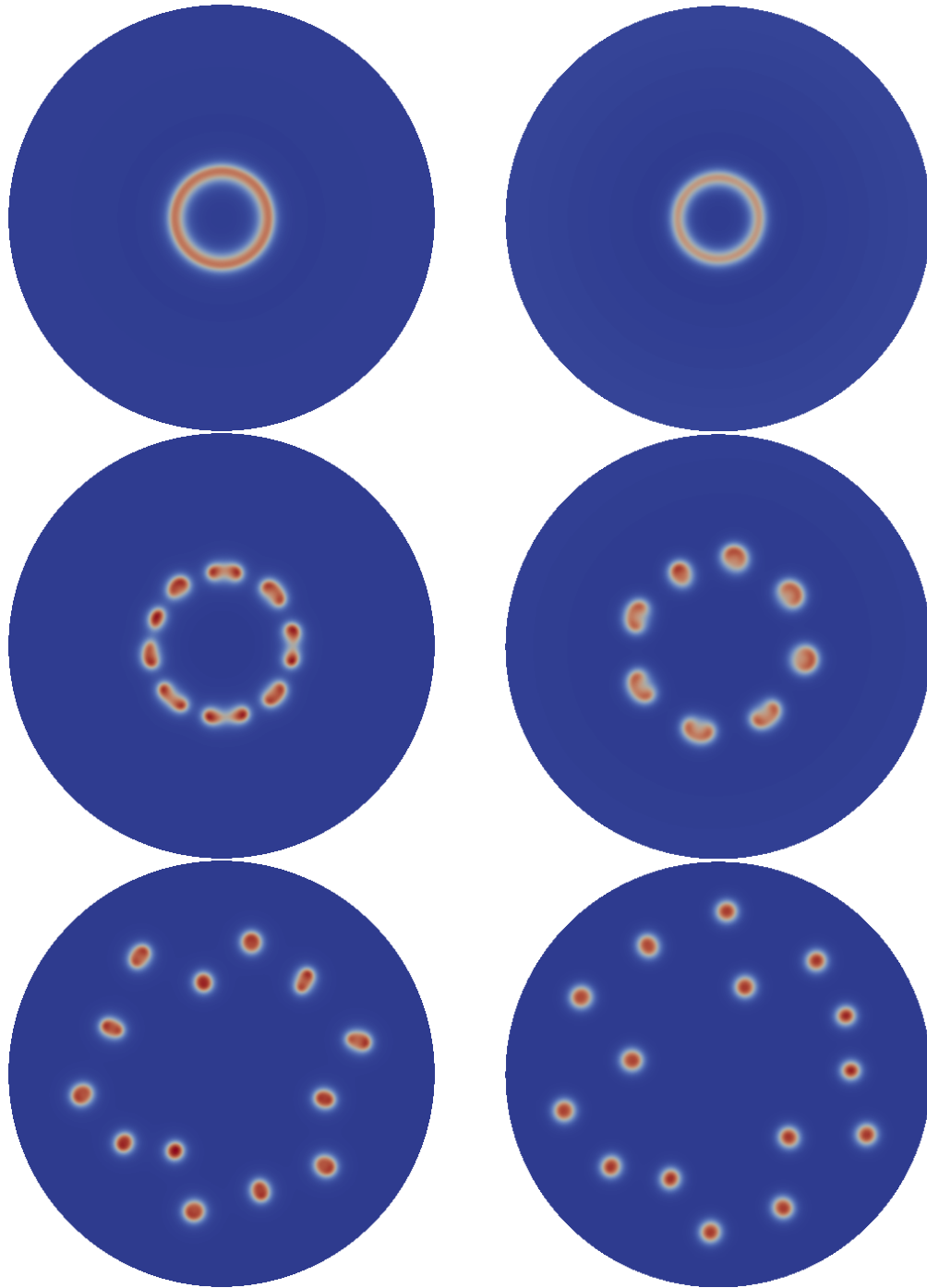


Figure 12.18: Localized spot solutions of the Brusselator system with $\alpha = \beta = 0.625$ (left) and $\alpha = \beta = 0.75$ (right). u is shown in both cases, and time progresses from top to bottom. The initial perturbation was identical in both cases. The initial perturbation in the centre of the domain forms a ring, whose radius is bigger if the fractional orders of diffusion α, β are smaller. The ring breaks up into several spots, which start to replicate and spread out over the whole domain. $n \approx 50,000$ unknowns were used in the finite element approximation.

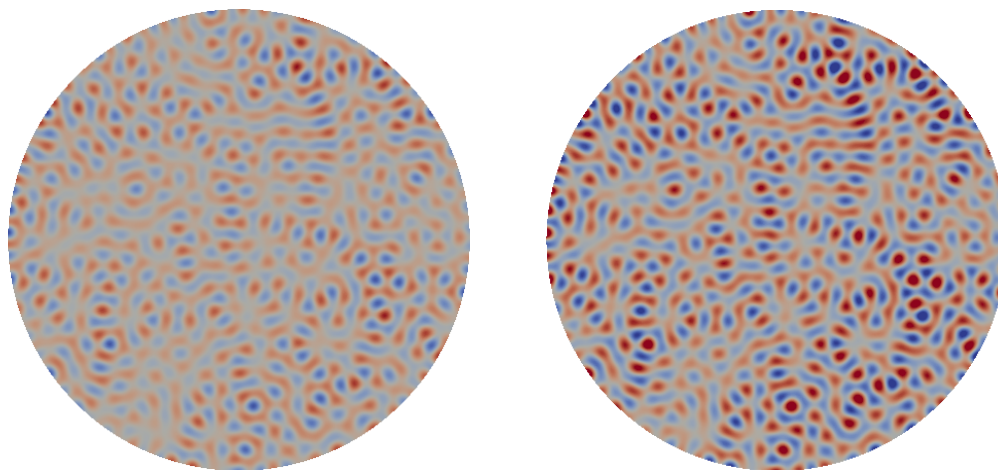


Figure 12.19: Stripe solutions of the Brusselator system with $\alpha = \beta = 0.75$. u is shown on the left, and v on the right. The random initial condition leads to the formation of stripes throughout the domain. $n \approx 50,000$ unknowns were used in the finite element approximation.

two-dimensional examples, we illustrated that the predicted optimal rates of convergence are obtained. We demonstrated how symmetries and periodicity in the problem can be leveraged to reduce the problem size even further.

Chapter 13

Future Work on Fractional and Non-Local Equations

In the present work, we focused on the case of fractional equations posed in $d = 1$ and $d = 2$ dimensions. The generalization to higher dimensions does not pose any fundamental difficulties. Expressions for the contributions of touching pairs of elements that contain a singularity would have to be rederived. Alternatively, Chernov, von Petersdorff, and Schwab proposed a procedure to transform quadrature rules with suitable weights from the hypercube $[0, 1]^{2d}$ to pairs of simplices [33, 34, 35]. The advantage of this method is that simplices in any dimension d can be handled. A preliminary implementation seems to suggest, however, that comparatively more quadrature points are necessary to obtain the same quadrature error.

Moreover, the use of higher order elements might be worth exploring. Preliminary results using piecewise quadratic instead of piecewise linear finite elements suggests that faster convergence can indeed be achieved, but a more thorough comparison of computational costs would be in order, since, for example, the computation of the error indicators introduced in Section 10.4 becomes more expensive.

At present, we also limited ourselves to homogeneous boundary conditions of Dirichlet and Neumann type. Non-homogeneous Dirichlet conditions can be taken

into account in a relatively straightforward way, provided that the contribution to the right-hand side of the exterior condition can be evaluated efficiently using numerical quadrature. The case of the Neumann condition is more subtle, as several generalizations of the integer order non-homogeneous Neumann condition have been proposed in the literature [40].

The approach described to obtain a sparse approximation to the dense system matrix for the fractional Laplacian does not rely strongly on the form of the interaction kernel $k(\vec{x}, \vec{y}) = |\vec{x} - \vec{y}|^{-(d+2s)}$, and generalizations to different kernels such as the one used in peridynamics [85] or kernels that can model anisotropic behavior are therefore possible.

Finally, it would be of interest to compare solutions to the fractional Poisson problem involving the integral definition of the fractional Laplacian with solutions obtained using the spectral fractional Laplacian. Solutions to the later can be obtained by solving an extension problem in $d+1$ dimensions [22]. Nochetto, Otárola, and Salgado leveraged this and proposed a method that converges essentially with rate $\mathcal{O}(n^{-1/(d+1)})$, n being the number of unknowns, provided that the solution is sufficiently regular [31, 80]. Since the two fractional operators, albeit different [84], seem to enjoy very similar properties, the question arises whether a method can be devised to solve equations involving the spectral fractional Laplacian with error decaying at the same rate as its integral counterpart, namely $\mathcal{O}(n^{-1/d})$.

Bibliography

- [1] Gabriel Acosta, Francisco M. Bersetche, and Juan-Pablo Borthagaray. “A short FE implementation for a 2d homogeneous Dirichlet problem of a Fractional Laplacian”. In: *ArXiv e-prints* (Oct. 2016).
- [2] Gabriel Acosta and Juan Pablo Borthagaray. “A fractional Laplace equation: regularity of solutions and Finite Element approximations”. In: *ArXiv e-prints* (July 2015).
- [3] Emmanuel Agullo, Luc Giraud, Pablo Salas, and Mawussi Zounon. “Interpolation-Restart Strategies for Resilient Eigensolvers”. In: *SIAM Journal on Scientific Computing* 38.5 (2016), pp. C560–C583. DOI: 10.1137/15M1042115.
- [4] Emmanuel Agullo, Luc Giraud, Pablo Salas, and Mawussi Zounon. “Recover-Restart Strategies for Resilient Parallel Numerical Linear Algebra Solvers”. In: *International Workshop on Parallel Matrix Algorithms and Applications (PMAA 2014)*. Lugano, Switzerland, July 2014.
- [5] Mark Ainsworth and Christian Glusa. “Aspects of an adaptive finite element method for the fractional Laplacian: a priori and a posteriori error estimates, efficient implementation and multigrid solver”.
- [6] Mark Ainsworth and Christian Glusa. “Is the multigrid method fault tolerant? The Multilevel Case.” In: *SIAM Journal on Scientific Computing* (Submitted).

- [7] Mark Ainsworth and Christian Glusa. “Is the Multigrid Method Fault Tolerant? The Two-Grid Case”. In: *SIAM Journal on Scientific Computing* 39.2 (2017), pp. C116–C143. DOI: 10.1137/16M1100691.
- [8] Mark Ainsworth and Christian Glusa. “Numerical Mathematics and Advanced Applications - ENUMATH 2015: Proceedings of ENUMATH 2015”. In: Cham: Springer International Publishing, 2015. Chap. Multigrid at Scale?, pp. 237–253. ISBN: 978-3-319-39929-4. DOI: 10.1007/978-3-319-39929-4_24.
- [9] Mark Ainsworth and Christian Glusa. “Towards an Efficient Finite Element Method for the Integral Fractional Laplacian on Polygonal Domains”. In: *TBD* (2017).
- [10] Mark Ainsworth, William McLean, and Thanh Tran. “The conditioning of boundary element equations on locally refined meshes and preconditioning by diagonal scaling”. In: *SIAM Journal on Numerical Analysis* 36.6 (1999), pp. 1901–1932.
- [11] Algirdas Avižienis, Jean-Claude Laprie, Brian Randell, and Carl Landwehr. “Basic concepts and taxonomy of dependable and secure computing”. In: *IEEE Transactions on Dependable and Secure Computing* 1.1 (2004), pp. 11–33.
- [12] Randolph E. Bank, Andrew H. Sherman, and Alan Weiser. “Some refinement algorithms and data structures for regular local mesh refinement”. In: *Scientific Computing, Applications of Mathematics and Computing to the Physical Sciences* 1 (1983), pp. 3–17.
- [13] Peter W. Bates. “On some nonlocal evolution equations arising in materials science”. In: *Nonlinear dynamics and evolution equations* 48 (2006), pp. 13–52.

- [14] Austin R Benson, Sven Schmit, and Robert Schreiber. “Silent Error Detection in Numerical Time-stepping Schemes”. In: *Int. J. High Perform. Comput. Appl.* 29.4 (Nov. 2015), pp. 403–421. ISSN: 1094-3420. DOI: 10.1177/1094342014532297.
- [15] David A. Benson, Stephen W. Wheatcraft, and Mark M. Meerschaert. “Application of a fractional advection-dispersion equation”. In: *Water Resources Research* 36.6 (2000), pp. 1403–1412.
- [16] Dimitri P Bertsekas. “Distributed asynchronous computation of fixed points”. In: *Mathematical Programming* 27.1 (1983), pp. 107–120.
- [17] Krzysztof Bogdan, Krzysztof Burdzy, and Zhen-Qing Chen. “Censored stable processes”. In: *Probability theory and related fields* 127.1 (2003), pp. 89–152.
- [18] Juan Pablo Borthagaray, Leandro M. Del Pezzo, and Sandra Martinez. “Finite element approximation for the fractional eigenvalue problem”. In: *ArXiv e-prints* (Mar. 2016).
- [19] Philippe Bougerol and Jean Lacroix. *Products of random matrices with applications to Schrödinger operators*. Vol. 8. Progress in Probability and Statistics. Boston, MA: Birkhäuser Boston Inc., 1985, pp. xii+283. ISBN: 0-8176-3324-3. DOI: 10.1007/978-1-4684-9172-2.
- [20] Dietrich Braess. *Finite elements. Theory, fast solvers and applications in solid mechanics. Translated from German by Larry L. Schumaker. 3rd ed.* English. Cambridge: Cambridge University Press, 2007. DOI: 10.1017/CB09780511618635.
- [21] James H. Bramble. *Multigrid methods*. Vol. 294. CRC Press, 1993.
- [22] Luis Caffarelli and Luis Silvestre. “An extension problem related to the fractional Laplacian”. In: *Communications in Partial Differential Equations* 32.8 (2007), pp. 1245–1260.

- [23] Xiao-Chuan Cai, Maksymilian Dryja, and Marcus Sarkis. “Restricted additive Schwarz preconditioners with harmonic overlap for symmetric positive definite linear systems”. In: *SIAM Journal on Numerical Analysis* 41.4 (2003), pp. 1209–1231.
- [24] Xiao-Chuan Cai and Marcus Sarkis. “A restricted additive Schwarz preconditioner for general sparse linear systems”. In: *SIAM Journal on Scientific Computing* 21.2 (1999), pp. 792–797.
- [25] Xiao-Chuan Cai and Jun Zou. “Some observations on the l2 convergence of the additive Schwarz preconditioned GMRES method”. In: *Numerical linear algebra with applications* 9.5 (2002), pp. 379–397.
- [26] Jon Calhoun, Luke Olson, Marc Snir, and William D. Gropp. “Towards a More Fault Resilient Multigrid Solver”. In: *Proceedings of the Symposium on High Performance Computing*. HPC ’15. Alexandria, Virginia: Society for Computer Simulation International, 2015, pp. 1–8. ISBN: 978-1-5108-0101-1.
- [27] Franck Cappello. “Fault tolerance in petascale/exascale systems: Current knowledge, challenges and research opportunities”. In: *International Journal of High Performance Computing Applications* 23.3 (2009), pp. 212–226.
- [28] Franck Cappello, Al Geist, Bill Gropp, Laxmikant Kale, Bill Kramer, and Marc Snir. “Toward exascale resilience”. In: *International Journal of High Performance Computing Applications* 23.4 (2009), pp. 374–388.
- [29] Franck Cappello, Al Geist, William Gropp, Sanjay Kale, Bill Kramer, and Marc Snir. “Toward exascale resilience: 2014 update”. In: *Supercomputing frontiers and innovations* 1.1 (2014), pp. 5–28.
- [30] Marc Casas, Bronis R. de Supinski, Greg Bronevetsky, and Martin Schulz. “Fault Resilience of the Algebraic Multi-grid Solver”. In: *Proceedings of the 26th ACM International Conference on Supercomputing*. ICS ’12. San Servolo

- Island, Venice, Italy: ACM, 2012, pp. 91–100. ISBN: 978-1-4503-1316-2. DOI: 10.1145/2304576.2304590.
- [31] Long Chen, Ricardo Nochetto, Enrique Otárola, and Abner Salgado. “Multi-level methods for nonuniformly elliptic operators and fractional diffusion”. In: *Mathematics of Computation* (2016).
- [32] Zhen-Qing Chen and Panki Kim. “Green function estimate for censored stable processes”. In: *Probability Theory and Related Fields* 124.4 (2002), pp. 595–610.
- [33] Alexey Chernov and Christoph Schwab. “Exponential convergence of Gauss–Jacobi quadratures for singular integrals over simplices in arbitrary dimension”. In: *SIAM Journal on Numerical Analysis* 50.3 (2012), pp. 1433–1455.
- [34] Alexey Chernov, Tobias von Petersdorff, and Christoph Schwab. “Exponential convergence of hp quadrature for integral operators with Gevrey kernels”. In: *ESAIM: Mathematical Modelling and Numerical Analysis* 45.3 (2011), pp. 387–422.
- [35] Alexey Chernov, Tobias von Petersdorff, and Christoph Schwab. “Quadrature algorithms for high dimensional singular integrands on simplices”. In: *Numerical Algorithms* 70.4 (2015), pp. 847–874.
- [36] Patrick Ciarlet. “Analysis of the Scott–Zhang interpolation in the fractional order Sobolev spaces”. In: *Journal of Numerical Mathematics* 21.3 (2013), pp. 173–180.
- [37] Andrea Crisanti, Giovanni Paladin, and Angelo Vulpiani. *Products of random matrices*. Springer, 1993.
- [38] Tao Cui, Jinchao Xu, and Chen-Song Zhang. “An error-resilient redundant subspace correction method”. In: *Computing and Visualization in Science* 18.2 (2017), pp. 65–77. ISSN: 1433-0369. DOI: 10.1007/s00791-016-0270-6.

- [39] Marta D’Elia and Max Gunzburger. “The fractional Laplacian operator on bounded domains as a special case of the nonlocal diffusion operator”. In: *Computers & Mathematics with Applications* 66.7 (2013), pp. 1245–1260. ISSN: 0898-1221. DOI: <http://dx.doi.org/10.1016/j.camwa.2013.07.022>.
- [40] Serena Dipierro, Xavier Ros-Oton, and Enrico Valdinoci. “Nonlocal problems with Neumann boundary conditions”. In: *ArXiv e-prints* (July 2014).
- [41] Victorita Dolean, Pierre Jolivet, and Frédéric Nataf. “An Introduction to Domain Decomposition Methods: algorithms, theory and parallel implementation”. France, Jan. 2015.
- [42] Jack Dongarra, Jeffrey Hittinger, John Bell, Luis Chacon, Robert Falgout, Michael Heroux, Paul Hovland, Esmond Ng, Clayton Webster, and Stefan Wild. *Applied Mathematics Research for Exascale Computing*. Tech. rep. Lawrence Livermore National Laboratory (LLNL), Livermore, CA, Feb. 2014. DOI: 10.2172/1149042.
- [43] Willy Dörfler. “A convergent adaptive algorithm for Poisson’s equation”. In: *SIAM Journal on Numerical Analysis* 33.3 (1996), pp. 1106–1124.
- [44] Michael G Duffy. “Quadrature over a pyramid or cube of integrands with a singularity at a vertex”. In: *SIAM journal on Numerical Analysis* 19.6 (1982), pp. 1260–1262.
- [45] Bartłomiej Dyda, Alexey Kuznetsov, and Mateusz Kwaśnicki. “Fractional Laplace Operator and Meijer G-function”. In: *Constructive Approximation* (2016), pp. 1–22. ISSN: 1432-0940. DOI: 10.1007/s00365-016-9336-4.
- [46] James Elliott, Mark Hoemmen, and Frank Mueller. “Evaluating the impact of SDC on the GMRES iterative solver”. In: *Parallel and Distributed Processing Symposium, 2014 IEEE 28th International*. IEEE. May 2014, pp. 1193–1202. DOI: 10.1109/IPDPS.2014.123.

- [47] James Elliott, Mark Hoemmen, and Frank Mueller. “Exploiting data representation for fault tolerance”. In: *Journal of Computational Science* (2016).
- [48] James Elliott, Kishor Kharbas, David Fiala, Frank Mueller, Kurt Ferreira, and Christian Engelmann. “Combining Partial Redundancy and Checkpointing for HPC”. In: *Proceedings of the 2012 IEEE 32Nd International Conference on Distributed Computing Systems*. ICDCS ’12. Washington, DC, USA: IEEE Computer Society, 2012, pp. 615–626. ISBN: 978-0-7695-4685-8. DOI: 10.1109/ICDCS.2012.56.
- [49] Mark Embree and Lloyd N. Trefethen. “Growth and Decay of Random Fibonacci Sequences”. English. In: *Proceedings: Mathematical, Physical and Engineering Sciences* 455.1987 (1999), pp. 2471–2485. ISSN: 13645021.
- [50] Alexandre Ern and Jean-Luc Guermond. *Theory and Practice of Finite Elements*. Applied Mathematical Sciences 159. New York, NY: Springer, 2004. DOI: 10.1007/978-1-4757-4355-5.
- [51] Birgit Faermann. “Localization of the Aronszajn-Slobodeckij norm and application to adaptive boundary element methods Part II. The three-dimensional case”. In: *Numerische Mathematik* 92.3 (2002), pp. 467–499.
- [52] Birgit Faermann. “Localization of the Aronszajn-Slobodeckij norm and application to adaptive boundary elements methods. Part I. The two-dimensional case”. In: *IMA Journal of Numerical Analysis* 20.2 (2000), pp. 203–234.
- [53] Kurt Ferreira, Jon Stearley, James H. Laros III, Ron Oldfield, Kevin Pedretti, Ron Brightwell, Rolf Riesen, Patrick G. Bridges, and Dorian Arnold. “Evaluating the Viability of Process Replication Reliability for Exascale Systems”. In: *Proceedings of 2011 International Conference for High Performance Computing, Networking, Storage and Analysis*. SC ’11. Seattle, Washington: ACM, 2011, 44:1–44:12. ISBN: 978-1-4503-0771-0. DOI: 10.1145/2063384.2063443.

- [54] Andreas Frommer and Daniel B Szyld. “On asynchronous iterations”. In: *Journal of Computational and Applied Mathematics* 123.1 (2000), pp. 201–216.
- [55] H. Furstenberg and H. Kesten. “Products of Random Matrices”. In: *The Annals of Mathematical Statistics* 31.2 (June 1960), pp. 457–469. DOI: 10 . 1214/aoms/1177705909.
- [56] Marc Gamell, Daniel S. Katz, Hemanth Kolla, Jacqueline Chen, Scott Klasky, and Manish Parashar. “Exploring Automatic, Online Failure Recovery for Scientific Applications at Extreme Scales”. In: *Proceedings of the International Conference for High Performance Computing, Networking, Storage and Analysis*. SC ’14. New Orleans, Louisiana: IEEE Press, 2014, pp. 895–906. ISBN: 978-1-4799-5500-8. DOI: 10 . 1109/SC.2014.78.
- [57] Paolo Gatto and Jan S Hesthaven. “Numerical approximation of the fractional laplacian via hp-finite elements, with an application to image denoising”. In: *Journal of Scientific Computing* 65.1 (2015), pp. 249–270.
- [58] R. K. Gettoor. “First passage times for symmetric stable processes in space”. In: *Transactions of the American Mathematical Society* 101.1 (1961), pp. 75–90.
- [59] Dominik Göddeke, Mirco Altenbernd, and Dirk Ribbrock. “Fault-tolerant finite-element multigrid algorithms with hierarchically compressed asynchronous checkpointing”. In: *Parallel Computing* 49 (2015), pp. 117–135.
- [60] Alexander A. Golovin, Bernard J. Matkowsky, and Vladimir A. Volpert. “Turing pattern formation in the Brusselator model with superdiffusion”. In: *SIAM Journal on Applied Mathematics* 69.1 (2008), pp. 251–272.
- [61] Gerd Grubb. “Fractional Laplacians on domains, a development of Hörmander’s theory of μ -transmission pseudodifferential operators”. In: *Advances in Mathematics* 268 (2015), pp. 478–528.

- [62] Wolfgang Hackbusch. *Iterative solution of large sparse systems of equations*. Vol. 95. Applied Mathematical Sciences. New York: Springer-Verlag, 1994, pp. xxii+429. ISBN: 0-387-94064-2. DOI: 10.1007/978-1-4612-4288-8.
- [63] Wolfgang Hackbusch. *Multi-grid methods and applications*. Vol. 4. Springer-Verlag Berlin, 1985. DOI: 10.1007/978-3-662-02427-0.
- [64] Thomas Herault and Yves Robert. *Fault-Tolerance Techniques for High-Performance Computing*. Springer International Publishing, 2015. DOI: 10.1007/978-3-319-20943-2.
- [65] Mark Hoemmen and Michael A Heroux. “Fault-tolerant iterative methods via selective reliability”. In: *Proceedings of the 2011 International Conference for High Performance Computing, Networking, Storage and Analysis (SC)*. IEEE Computer Society. Vol. 3. 2011, p. 9.
- [66] Roger A. Horn and Charles R. Johnson. *Matrix analysis*. Cambridge University press, 2012.
- [67] Roger A. Horn and Charles R. Johnson. *Topics in matrix analysis*. Cambridge University press, 1991. DOI: 10.1017/cbo9780511840371.
- [68] Kuang-Hua Huang and Jacob Abraham. “Algorithm-based fault tolerance for matrix operations”. In: *Computers, IEEE Transactions on* 100.6 (1984), pp. 518–528.
- [69] Markus Huber, Björn Gmeiner, Ulrich Rüde, and Barbara Wohlmuth. “Resilience for Multigrid Software at the Extreme Scale”. In: *arXiv preprint arXiv:1506.06185* (2015).
- [70] Luc Jaulmes, Marc Casas, Miquel Moretó, Eduard Ayguadé, Jesús Labarta, and Mateo Valero. “Exploiting Asynchrony from Exact Forward Recovery for DUE in Iterative Solvers”. In: *Proceedings of the International Conference for High Performance Computing, Networking, Storage and Analysis*. SC ’15.

- Austin, Texas: ACM, 2015, 53:1–53:12. ISBN: 978-1-4503-3723-6. DOI: 10.1145/2807591.2807599.
- [71] George Karypis and Vipin Kumar. “A Fast and High Quality Multilevel Scheme for Partitioning Irregular Graphs”. In: *SIAM Journal on Scientific Computing* 20.1 (1998), pp. 359–392. DOI: 10.1137/S1064827595287997.
- [72] Toshiyuki Koto. “IMEX Runge–Kutta schemes for reaction–diffusion equations”. In: *Journal of Computational and Applied Mathematics* 215.1 (2008), pp. 182–195.
- [73] Tim Kröger and Tobias Preusser. “Stability of the 8-tetrahedra shortest-interior-edge partitioning method”. In: *Numerische Mathematik* 109.3 (2008), pp. 435–457. ISSN: 0945-3245. DOI: 10.1007/s00211-008-0148-8.
- [74] Shuangzhe Liu and Götz Trenkler. “Hadamard, Khatri-Rao, Kronecker and other matrix products”. In: *Int. J. Inform. Syst. Sci* 4.1 (2008), pp. 160–177.
- [75] Yifei Lou, Xiaoqun Zhang, Stanley Osher, and Andrea Bertozzi. “Image recovery via nonlocal operators”. In: *Journal of Scientific Computing* 42.2 (2010), pp. 185–197.
- [76] Frédéric Magoulès, Daniel B. Szyld, and Cédric Venet. “Asynchronous optimized Schwarz methods with and without overlap”. In: *Numerische Mathematik* (2017), pp. 1–29. ISSN: 0945-3245. DOI: 10.1007/s00211-017-0872-z.
- [77] B. M. McCay and M. N. L. Narasimhan. “Theory of nonlocal electromagnetic fluids”. In: *Archives of Mechanics* 33.3 (1981), pp. 365–384.
- [78] William Charles Hector McLean. *Strongly Elliptic Systems and Boundary Integral Equations*. Cambridge University Press, 2000.
- [79] Mark M. Meerschaert and Alla Sikorskii. *Stochastic models for fractional calculus*. Vol. 43. de Gruyter Studies in Mathematics. Walter de Gruyter & Co., Berlin, 2012, pp. x+291. ISBN: 978-3-11-025869-1.

- [80] Ricardo H Nochetto, Enrique Otárola, and Abner J Salgado. “A PDE approach to fractional diffusion in general domains: a priori error analysis”. In: *Foundations of Computational Mathematics* 15.3 (2015), pp. 733–791.
- [81] Piyush Sao and Richard Vuduc. “Self-stabilizing Iterative Solvers”. In: *Proceedings of the Workshop on Latest Advances in Scalable Algorithms for Large-Scale Systems*. ScalA ’13. Denver, Colorado: ACM, 2013, 4:1–4:8. ISBN: 978-1-4503-2508-0. DOI: 10.1145/2530268.2530272.
- [82] Stefan A. Sauter and Christoph Schwab. “Boundary element methods”. In: *Boundary Element Methods*. Springer, 2010, pp. 183–287.
- [83] L Ridgway Scott and Shangyou Zhang. “Finite element interpolation of non-smooth functions satisfying boundary conditions”. In: *Mathematics of Computation* 54.190 (1990), pp. 483–493.
- [84] Raffaella Servadei and Enrico Valdinoci. “On the spectrum of two different fractional operators”. In: *Proceedings of the Royal Society of Edinburgh: Section A Mathematics* 144.04 (2014), pp. 831–855.
- [85] Stewart A. Silling. “Reformulation of elasticity theory for discontinuities and long-range forces”. In: *Journal of the Mechanics and Physics of Solids* 48.1 (2000), pp. 175–209.
- [86] Ian H. Sloan. “Error analysis of boundary integral methods”. In: *Acta Numerica* 1 (1992), pp. 287–339.
- [87] Ian H. Sloan and Alastair Spence. “The Galerkin method for integral equations of the first kind with logarithmic kernel: Theory”. In: *IMA Journal of Numerical Analysis* 8.1 (1988), pp. 105–122.
- [88] Barry Smith, Petter Bjorstad, and William Gropp. *Domain decomposition: parallel multilevel methods for elliptic partial differential equations*. Cambridge university press, 2004.

- [89] Marc Snir, Robert W Wisniewski, Jacob A Abraham, Sarita V Adve, Saurabh Bagchi, Pavan Balaji, Jim Belak, Pradip Bose, Franck Cappello, Bill Carlson, et al. “Addressing failures in exascale computing”. In: *International Journal of High Performance Computing Applications* 28.2 (2014), pp. 129–173.
- [90] Miroslav Stoyanov and Clayton Webster. “Numerical Analysis of Fixed Point Algorithms in the Presence of Hardware Faults”. In: *SIAM Journal on Scientific Computing* 37.5 (2015), pp. C532–C553. DOI: 10.1137/140991406.
- [91] John C. Strikwerda. “A probabilistic analysis of asynchronous iteration”. In: *Linear Algebra Appl.* 349 (2002), pp. 125–154. ISSN: 0024-3795. DOI: 10.1016/S0024-3795(02)00258-6.
- [92] Arthur H. Stroud. *Approximate calculation of multiple integrals*. Prentice-Hall, 1971.
- [93] Peter Tankov. *Financial modelling with jump processes*. Vol. 2. CRC press, 2003.
- [94] Andrea Toselli and Olof Widlund. *Domain decomposition methods: algorithms and theory*. Vol. 3. Springer, 2005.
- [95] Ulrich Trottenberg, Cornelius W. Oosterlee, and Anton Schüller. *Multigrid*. San Diego, CA: Academic Press Inc., 2001, pp. xvi+631. ISBN: 0-12-701070-X.
- [96] Enrico Valdinoci. “From the long jump random walk to the fractional Laplacian”. In: *SeMA Journal: Boletín de la Sociedad Española de Matemática Aplicada* 49 (2009), pp. 33–44.
- [97] Jacques Vanneste. “Estimating generalized Lyapunov exponents for products of random matrices”. In: *Physical Review E* 81.3 (2010), p. 036701.
- [98] Divakar Viswanath. “Random Fibonacci sequences and the number 1.13198824...”. In: *Math. Comp.* 69.231 (2000), pp. 1131–1155. ISSN: 0025-5718. DOI: 10.1090/S0025-5718-99-01145-X.

- [99] Eric W. Weisstein. *Meijer G-Function*. From *MathWorld—A Wolfram Web Resource*. 2017.
- [100] Bruce J. West. *Fractional Calculus View of Complexity: Tomorrow's Science*. CRC Press, 2016.
- [101] Jinchao Xu. *An introduction to multilevel methods*. 1997.
- [102] Jinchao Xu. “Iterative methods by space decomposition and subspace correction”. In: *SIAM Rev.* 34.4 (1992), pp. 581–613. ISSN: 0036-1445. DOI: 10.1137/1034116.
- [103] Yeli Yan, Ian H. Sloan, et al. *On integral equations of the first kind with logarithmic kernels*. University of NSW, 1988.