



AFRL-AFOSR-VA-TR-2019-0220

Optimal and unstructured high-order non-intrusive approximations for uncertain parameterized simulations

Akil Narayan
UNIVERSITY OF UTAH SALT LAKE CITY

07/11/2019
Final Report

DISTRIBUTION A: Distribution approved for public release.

Air Force Research Laboratory
AF Office Of Scientific Research (AFOSR)/ RTA2
Arlington, Virginia 22203
Air Force Materiel Command

REPORT DOCUMENTATION PAGE					Form Approved OMB No. 0704-0188	
The public reporting burden for this collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing the burden, to Department of Defense, Washington Headquarters Services, Directorate for Information Operations and Reports (0704-0188), 1215 Jefferson Davis Highway, Suite 1204, Arlington, VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to any penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number. PLEASE DO NOT RETURN YOUR FORM TO THE ABOVE ADDRESS.						
1. REPORT DATE (DD-MM-YYYY)		2. REPORT TYPE			3. DATES COVERED (From - To)	
4. TITLE AND SUBTITLE				5a. CONTRACT NUMBER		
				5b. GRANT NUMBER		
				5c. PROGRAM ELEMENT NUMBER		
6. AUTHOR(S) Narayan, Akil				5d. PROJECT NUMBER		
				5e. TASK NUMBER		
				5f. WORK UNIT NUMBER		
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES)					8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES) Air Force Office of Scientific Research 875 N. Randolph, Ste.325 Arlington Virginia, 22203					10. SPONSOR/MONITOR'S ACRONYM(S) AFOSR	
					11. SPONSOR/MONITOR'S REPORT NUMBER(S)	
12. DISTRIBUTION/AVAILABILITY STATEMENT Distribution A -- Approved for public release; distribution unlimited						
13. SUPPLEMENTARY NOTES						
14. ABSTRACT This is the final report for AFOSR project FA9550-15-1-0467. This report summarizes the content of research and publications completed under this project. 23 published articles and an additional 6 in-review articles were produced by this project. The main project goals are design of novel mathematical algorithms for constructing sample-based approximation in						
15. SUBJECT TERMS Sample-based approximation, least squares, compressed sensing, interpolation, model reduction, high-dimensional approximation, approximation optimality						
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT UU	18. NUMBER OF PAGES 21	19a. NAME OF RESPONSIBLE PERSON Narayan, Akil	
a. REPORT U	b. ABSTRACT U	c. THIS PAGE U			19b. TELEPHONE NUMBER (Include area code) 765-444-9391	

Final Report

“Optimal and unstructured high-order non-intrusive approximations for uncertain parameterized simulations”

PI: Akil Narayan
The University of Utah

AFOSR award FA9550-15-1-0467
Reporting period: Sept 30, 2015 - March 29, 2019

Contents

1	Executive summary – outline, objectives, accomplishments	3
2	Project personnel and publications	5
2.1	Personnel	5
2.2	Publications	5
3	Technical description of accomplishments	8
3.1	Sampling for parametric least-squares and compressive sampling	8
3.1.1	CS with quadrature subsampling	9
3.1.2	CS with Christoffel Sparse Approximation (CSA)	10
3.1.3	Optimal quadrature rules: characterization and computation	12
3.2	Application: sampling in multifidelity approximation	12
3.2.1	Convergence acceleration of multifidelity models	15
3.3	Application: Sampling for reduced order modeling	16
3.3.1	Model reduction for nonlocal equations	18
3.4	Application: Topology optimization	18
3.4.1	Machine learning samplers	19
3.4.2	Design under uncertainty	19

1 Executive summary – outline, objectives, accomplishments

The approximation and prediction of output quantities of interest in large-scale simulation software is an ongoing challenge in scientific computing. This difficulty is compounded when the simulation software contains numerous tunable input parameters that specify modeling scenarios, geometry, and uncertainty. The main goal of this project is robust and efficient prediction of the variability of quantities of interest with respect to these input parameters. This is primarily accomplished via non-intrusive sampling of models. Straightforward and naive sampling methods often (usually) yield suboptimal performance and convergence guarantees. **This project aims to develop novel, modern sampling strategies that perform well and are provably convergent, ideally without dependence on dimension. Nearing the end of this project, the efficacy of the developed procedures will be tested on realistic parameterized scientific problems.**

Let a function $f(x)$, $f : D \rightarrow \mathbb{R}$ represent the scalar quantity of interest, where $D \subset \mathbb{R}^d$ is a d -dimensional parameter space. We seek approximations with N degrees of freedom of the form

$$f(x) \simeq f_N(x) = \sum_{n=1}^N \hat{f}_n \phi_n(x), \quad (1)$$

where the ϕ_n are basis functions that encode the variability of f with respect to x . Given $f(x_j)$ for $x_j \in \mathbb{R}^d$ and $j = 1, \dots, M$, we seek to compute the coefficients $\hat{f}_n$ in as stable and accurate a manner as possible. Accuracy is measured in a weighted norm on $D \subset \mathbb{R}^d$, where the weight function is $w(x)$. Since M is a surrogate for the amount of work required to compute this approximation (we must sample $f(x)$ M times), the goal is to minimize M with respect to the number of degrees of freedom in the approximation, N .

In this report we refer to the problem of designing samples to maximize stability and accuracy as the problem of *optimal approximation*. The problem of optimal approximation is an open problem for $d > 1$. This project considers both cases when the ϕ_n are specified *a priori* (e.g., polynomials) and when they are learned from simulation data (e.g., reduced order modeling).

This project seeks to construct optimal approximations based on mathematically rigorous properties of sampling. This is accomplished via following **overall project tasks**:

- Construction of near-optimal randomized and deterministic sampling sequences for non-adaptive approximation of the form (1). This requires developing collocation-based quadrature, interpolation, discrete least-squares, and compressive sampling schemes to recover the coefficients $\hat{f}_n$ in (1).
- Development of adaptive approximation strategies, where the functions ϕ_n in (1) are constructed on-the-fly. The main goals in the context of this project are such adaptive strategies that utilize models of different fidelities and accuracies that synthesize a superior prediction from the individual models.
- Application to parameterized problems in scientific computing. The terminal application of this project seeks to understand the effects that tunable parameters have on output quantities of interest.

The **deliverables** for this project are the following:

- Publication or acceptance of 23 articles: [P19, P5, P9, P7, P22, P2, P12, P13, P24, P1, P18, P11, P14, P23, P8, P3, P15, P21, P6, P20, P10, P4, P16].
- 6 References [S4, S1, S6, S2, S3, S5] are submitted the time of writing of this report. All but one are in the revision stage and are expected to be accepted for publication.

The **accomplishments** for this project are the following:

- Development sampling strategies and convergence guarantees for novel kinds of sample-based approximation techniques: quadrature [P14, P11], [S4], least squares [P19, P5, P7, P22, P18], [S3] and compressive sampling [P7, P6, P1, P9].
- Dissemination of software that accomplishes near-optimal randomized sampling [P17].
- Design of novel sampling techniques for adaptive surrogate construction, applied to reduced order modeling and multi-fidelity modeling [P8, P3, P21, P20].
- Convergence analysis for adaptive techniques [P3, P15].

Applications:

- Structure design in topology optimization [P14, P16], [S5].
- High-order PDE solvers on manifolds [P23].
- Parametric reduced order modeling with multifidelity techniques [P8, P20]
- Reduced order modeling for nonlocal/fractional diffusion equations [S1, S2].
- Time-dependent data assimilation [P24], [S6]

2 Project personnel and publications

2.1 Personnel

- Akil Narayan, PI, Assistant Professor, University of Utah
- Vahid Keshavarzzadeh, Postdoctoral scholar, University of Utah
- Mani Razi, Postdoctoral scholar, University of Utah

2.2 Publications

Articles submitted

- [S1] H. ANTIL, Y. CHEN, AND A. NARAYAN, *Certified reduced basis methods for fractional Laplace equations via extension*, arXiv:1808.00584 [math], (2018). arXiv: 1808.00584.
- [S2] H. DINH, H. ANTIL, Y. CHEN, E. CHERKAEV, AND A. NARAYAN, *Model reduction for fractional elliptic problems using Kato's formula*, arXiv:1904.09332 [math], (2019). arXiv: 1904.09332.
- [S3] L. GUO, A. NARAYAN, AND T. ZHOU, *Constructing least squares polynomial approximations*, submitted, (2019).
- [S4] V. KESHAVARZZADEH, R. M. KIRBY, AND A. NARAYAN, *Generation of Nested Quadrature Rules for Generic Weight Functions via Numerical Optimization: Application to Sparse Grids*, arXiv:1808.03707 [math], (2018). arXiv: 1808.03707.
- [S5] ———, *Stress-based topology optimization under uncertainty via simulation-based gaussian process*, Submitted, (2019).
- [S6] L. YANG, Y. QIN, A. NARAYAN, AND P. WANG, *Data assimilation for models with parametric uncertainty*, Submitted, (2019).

Articles published and/or accepted during award period

- [P1] B. ADCOCK, A. BAO, J. JAKEMAN, AND A. NARAYAN, *Compressed sensing with sparse corruptions: Fault-tolerant sparse collocation approximations*, SIAM Journal on Uncertainty Quantification (to appear), (2017). arXiv: 1703.00135 [math.NA].
- [P2] L. BOS, A. NARAYAN, N. LEVENBERG, AND F. PIAZZON, *An Orthogonality Property of the Legendre Polynomials*, Constructive Approximation, 45 (2017), pp. 65–81. arXiv:1505.06635 [math.CA].
- [P3] Y. CHEN, J. JIANG, AND A. NARAYAN, *A robust error estimator and a residual-free error indicator for reduced basis methods*, Computers & Mathematics with Applications, 77 (2019), pp. 1963–1979. arxiv: 1710.08999 [math.NA].

- [P4] M. CHENG, A. NARAYAN, Y. QIN, P. WANG, X. ZHONG, AND X. ZHU, *An efficient solver for cumulative density function-based solutions of uncertain kinematic wave models*, Journal of Computational Physics, 382 (2019), pp. 138–151. arXiv:1901.08520 [math.NA].
- [P5] L. GUO, A. NARAYAN, L. YAN, AND T. ZHOU, *Weighted Approximate Fekete Points: Sampling for Least-Squares Polynomial Approximation*, SIAM Journal on Scientific Computing, 40 (2018), pp. A366–A387. arXiv:1708.01296 [math.NA].
- [P6] L. GUO, A. NARAYAN, AND T. ZHOU, *A gradient enhanced ℓ_1 -minimization for sparse approximation of polynomial chaos expansions*, Journal of Computational Physics, 367 (2018), pp. 49–64.
- [P7] L. GUO, A. NARAYAN, T. ZHOU, AND Y. CHEN, *Stochastic Collocation Methods via L_1 Minimization Using Randomized Quadratures*, SIAM Journal on Scientific Computing, 39 (2017), pp. A333–A359. arXiv: 1602.00995 [math.NA].
- [P8] J. HAMPTON, H. R. FAIRBANKS, A. NARAYAN, AND A. DOOSTAN, *Practical error bounds for a non-intrusive bi-fidelity approach to parametric/stochastic model reduction*, Journal of Computational Physics, 368 (2018), pp. 315–332. arXiv: 1709.03661.
- [P9] J. JAKEMAN, A. NARAYAN, AND T. ZHOU, *A Generalized Sampling and Preconditioning Scheme for Sparse Approximation of Polynomial Chaos Expansions*, SIAM Journal on Scientific Computing, 39 (2017), pp. A1114–A1144. arXiv: 1602.06879 [math.NA].
- [P10] J. D. JAKEMAN, F. FRANZELIN, A. NARAYAN, M. ELDRED, AND D. PLFÜGER, *Polynomial chaos expansions for dependent random variables*, Computer Methods in Applied Mechanics and Engineering, 351 (2019), pp. 643–666.
- [P11] J. D. JAKEMAN AND A. NARAYAN, *Generation and application of multivariate polynomial quadrature rules*, Computer Methods in Applied Mechanics and Engineering, 338 (2018), pp. 134–161. arXiv:1711.00506 [math].
- [P12] J. JIANG, Y. CHEN, AND A. NARAYAN, *A Goal-Oriented Reduced Basis Methods-Accelerated Generalized Polynomial Chaos Algorithm*, SIAM/ASA Journal on Uncertainty Quantification (to appear), 4 (2016), pp. 1398–1420. arXiv:1601.00137 [math.NA].
- [P13] J. JIANG, Y. CHEN, AND A. NARAYAN, *Offline-Enhanced Reduced Basis Method through adaptive construction of the Surrogate Training Set*, Journal of Scientific Computing (to appear), (2017). arXiv: 1703.05683 [math].
- [P14] V. KESHAVARZZADEH, R. KIRBY, AND A. NARAYAN, *Numerical Integration in Multiple Dimensions with Designed Quadrature*, SIAM Journal on Scientific Computing, 40 (2018), pp. A2033–A2061. arXiv:1804.06501 [cs.NA].
- [P15] V. KESHAVARZZADEH, R. M. KIRBY, AND A. NARAYAN, *Convergence Acceleration for Time Dependent Parametric Multifidelity Models*, SIAM Journal on Numerical Analysis (to appear), arXiv:1808.03379 [math], (2018). arXiv: 1808.03379.
- [P16] ———, *Parametric Topology Optimization with Multi-Resolution Finite Element Models*, International Journal for Numerical Methods in Engineering, (2018). arXiv: 1808.10367 [cs].

- [P17] A. NARAYAN, *akilnarayan/induced-distributions: Initial beta release*, 2017. DOI: 10.5281/zenodo.569416.
- [P18] ———, *Computation of Induced Orthogonal Polynomial Distributions*, Electronic Transactions on Numerical Analysis, 50 (2017), pp. 71–97. arXiv:1704.08465 [math].
- [P19] A. NARAYAN, J. JAKEMAN, AND T. ZHOU, *A Christoffel function weighted least squares algorithm for collocation approximations*, Mathematics of Computation, 86 (2017), pp. 1913–1947. arXiv: 1412.4305 [math.NA].
- [P20] D. PERRY, R. KIRBY, A. NARAYAN, AND R. WHITAKER, *Allocation Strategies for High Fidelity Models in the Multifidelity Regime*, SIAM/ASA Journal on Uncertainty Quantification, 7 (2019), pp. 203–231. arXiv: 1812.11601 [math.NA].
- [P21] R. PULCH AND A. NARAYAN, *Balanced truncation for model order reduction of linear dynamical systems with quadratic outputs*, SIAM Journal on Scientific Computing (accepted), arXiv:1709.06677 [math], (2017). arXiv: 1709.06677.
- [P22] P. SESHADRI, A. NARAYAN, AND S. MAHADEVAN, *Optimal Quadrature Subsampling for Least Squares Polynomial Approximations*, SIAM/ASA Journal on Uncertainty Quantification (to appear), (2017). arXiv: 1601.05470 [math.NA].
- [P23] V. SHANKAR, A. NARAYAN, AND R. M. KIRBY, *RBF-LOI: Augmenting Radial Basis Functions (RBFs) with Least Orthogonal Interpolation (LOI) for solving PDEs on surfaces*, Journal of Computational Physics, 373 (2018), pp. 722–735. arXiv: 1807.02775.
- [P24] L. YANG, A. NARAYAN, AND P. WANG, *Sequential data assimilation with multiple nonlinear models and applications to subsurface flow*, Journal of Computational Physics, 346 (2017), pp. 356–368. arXiv: 1707.06394.

3 Technical description of accomplishments

3.1 Sampling for parametric least-squares and compressive sampling

The accuracy of computed coefficients $\hat{f}_n$ in (1) is usually measured in terms of best approximations. For example, if ϕ_n form an orthonormal basis, one may seek the best L^2 coefficients,

$$\hat{f}_n^* = \int_D f(x) \phi_n(x) d\mu(x), \quad f_N^* = \sum_{n=1}^N \hat{f}_n^* \phi_n(x)$$

where μ is a known probability measure on D . When performing overdetermined (large-data) least-squares approximations, the goal is establishment of the estimate

$$\|f_N - f\|_{L_\mu^2} \lesssim C \|f_N^* - f\|_{L_\mu^2},$$

where ideally the constant C is an absolute constant not dependent on dimension d . When performing underdetermined (small-data) approximations, one cannot hope to obtain all coefficients from incomplete data, so the goal is often reframed in terms of best s -term estimates. Such estimates take the form

$$\|f_N - f_N^*\|_{L_\mu^2} \lesssim C \inf_{\|c\|_0 \leq s} \|\hat{f}^* - \hat{c}\|_2,$$

and the machinery of compressive sampling yields such estimates with $m \sim s$ data points.

The major problem that this project seeks to overcome is that the constants C usually depend heavily on d , D , and the basis ϕ_n when data is obtained using standard sampling techniques. We focused on development of novel sampling techniques that overcome this restriction. Much work has shown that sampling data from the probability measure defined by

$$d\mu_N(x) = \frac{1}{N} \sum_{n=1}^N \phi_n^2(x) d\mu(x),$$

allows one to improve estimates significantly, almost eliminating the dependence of C on d , D , and ϕ_n . For example, one of the novel results in [P19] states that if the number of samples M drawn from the measure $d\mu_N$ satisfies $M \geq kN \log N$, then

$$\|f_N - f\|_{L_\mu^2} \leq \left(1 + \frac{1}{\log M}\right) \|f_N^* - f\|_{L_\mu^2} + \frac{1}{M^{k-1}},$$

for very general d , D , and ϕ_n . Thus, *log-linear* sampling of M with respect to N is enough to establish very general convergence guarantees in any dimension.

This work has been exploited for constructing optimal least-squares constructions in various settings [P19, P5], and for compressive sampling approximations [P1, P9].

In the following two sections we describe two particular techniques for compressed sampling: a subsampling method from [P7] and a weighting technique from [P9]. Following this we describe novel quadrature rule theory and computational constructions developed in this project in [P11, P14].

3.1.1 CS with quadrature subsampling

Given $n \in \mathbb{N}$, we can form the n^d -point tensor-product grid, formed from n univariate Gauss quadrature points in each dimension. We likewise form the product Gauss quadrature weights on this tensor-product grid. (Note that we do not actually form this d -dimensional tensorial grid explicitly, which would be too expensive.) We randomly subsample M nodes $x_1, \dots, x_N$ and weights $w_1^2, \dots, w_M^2$ from this grid. The nodes x_n are used to sample a function, and the weights w_n precondition the solver to make computations well-conditioned. This procedure is simple, straightforward, and with M fixed has only linear dependence on dimension.

We investigated this procedure in [P7], and there the following convergence result was established: let a sparsity level s be such that $M \gtrsim K(w)s$. Then with high probability the compressive sampling solution c satisfies

$$\|c - \tilde{f}\|_2 \leq \frac{C_1 \sigma_{s,1}(\tilde{f})}{\sqrt{s}} + C_2 \varepsilon, \quad (2)$$

where C_1 and C_2 are universal constants, and $\tilde{f}$ is the unknown vector of exact expansion coefficients of $f(x)$. The quantity $\sigma_{s,1}$ is the best s -term error in the 1-norm:

$$\sigma_{s,1}(v) = \min_{\|w\|_0 \leq s} \|v - w\|_1.$$

We show that in one dimension $K(w)$, which determines the requisite sample size M , has the following behavior

- (Beta distributions) $w(x) = (1-x)^\alpha(1+x)^\beta$ for $\alpha, \beta > -1$ on $x \in [-1, 1]$. Then $K(w) = C(\alpha, \beta)$ is a constant depending (linearly) on α and β .
- (Two-sided exponential distributions) $w(x) = \exp(-|x|^\alpha)$ for $\alpha > 1$ on $x \in \mathbb{R}$. Then

$$K(w) = \begin{cases} C(\alpha)k^{2/3}, & \alpha \geq \frac{3}{2} \\ C(\alpha)k^{1/\alpha}, & 1 < \alpha < \frac{3}{2} \end{cases}$$

where k is the maximum polynomial degree in the dictionary defined by the Vandermonde matrix V .

- (One-sided exponential distributions) $w(x) = \exp(-x^\alpha)$ for $\alpha > \frac{1}{2}$ on $x \in [0, \infty)$. Then

$$K(w) = \begin{cases} C(\alpha)k^{2/3}, & \alpha \geq \frac{3}{4} \\ C(\alpha)k^{1/2\alpha}, & \frac{1}{2} < \alpha < \frac{3}{4} \end{cases}$$

where k is the maximum polynomial degree in the dictionary defined by the Vandermonde matrix V .

The above behaviors of $K(w)$ are sharp: no smaller restriction on the sample count can prove convergence using the strategy of bounded orthonormal systems that we leverage. All of these results generalize to the multidimensional case via tensor-products. Note that our results are very general and cover nearly all types of continuous distributions one encounters in practice.

In terms of practicality of this procedure, we show summary results in Figure 1. Both plots show the superiority of quadrature subsampling compared to standard method using other iid sampling algorithms.

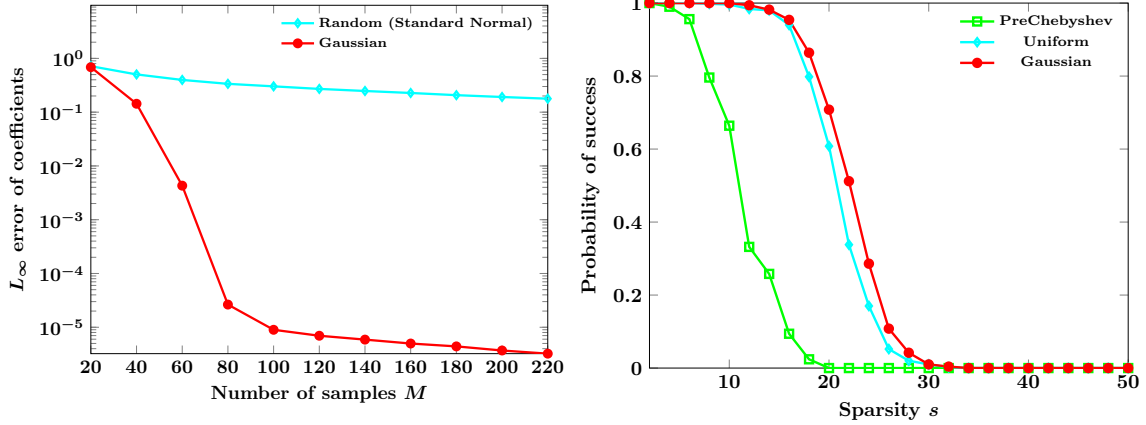


Figure 1: Results from [P7]. Recovery using random quadrature subsampling (“Gaussian”). Left: Error versus sample count for a smooth test function in $d = 2$ dimensions with $w(x)$ corresponding to a normal distribution. Right: Empirical probability of recovery of exactly sparse polynomials in $d = 10$ dimensions with w the uniform weight on $[-1, 1]^{10}$. In the left and right plots, comparison is shown against solving and ℓ_1 optimization problem using unweighted iid sampling according to w (“Random” and “Uniform”, respectively), and in the right plot, a comparison against a popular preconditioned Chebyshev sampling is also shown.

3.1.2 CS with Christoffel Sparse Approximation (CSA)

A second approach for sparse approximation via compressive sampling has been developed in [P9]. This method for sparse approximation utilizes results from pluripotential theory. This procedure solves a preconditioned ℓ^1 minimization problem by sampling x_n iid from a weighted pluripotential equilibrium measure $\mu_{D,w}$. The preconditioning weights correspond to ℓ^2 row-norms of the design matrix. These row-norms are evaluations of the Christoffel function, hence the “C” in CSA.

When using this procedure in one dimension and tensor-product situations, our convergence results are the same as in (2), with $K(w)$ behaving exactly as described in the previous section for those weights. However, the CSA method can be applied to *any* case when w is continuous, not just the tensor-product case. Thus, the method is more general than the quadrature subsampling strategy.

To illustrate the effectiveness of the CSA procedure, Figure 2 shows transition plots when compared again a standard Monte Carlo ℓ^1 recovery procedure in both low and high dimensions, on bounded and unbounded domains. The CSA method is a simple, general procedure for recovering sparse expansions.

In Figure 3 we solve a x -parameterized diffusion differential equation. When compared against standard Monte Carlo as well as preconditioned Chebyshev sampling, the CSA procedure produces superior results.

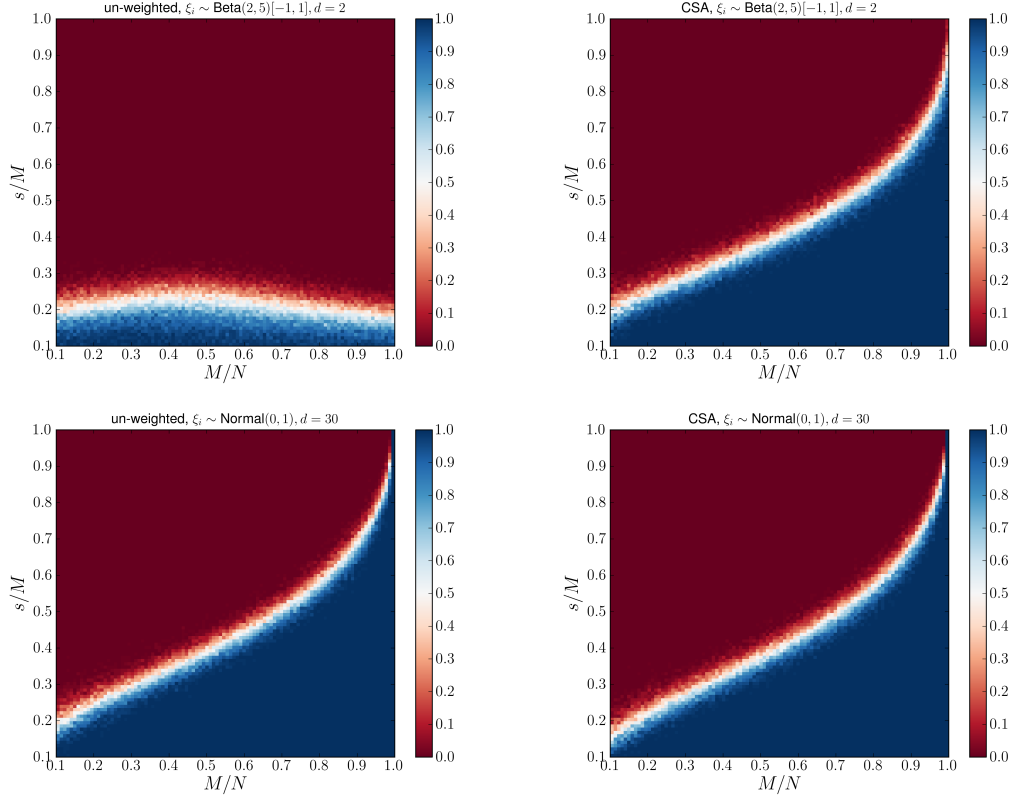


Figure 2: Results from [P9]. Compressive sampling transition plots for the CSA procedure. The horizontal axis is the sample count size relative to the size of the dictionary, and the vertical axis is the sparsity of a function relative to the sample count. Each pixel is colored with probability is successfully recovering the sparse solution from ℓ_1 minimization. Top row: $d = 2$ dimensions with w an asymmetric Beta density. Bottom row: $d = 30$ dimensions with w a normal random variable density. Left column: Unweighted ℓ_1 optimization solved by sampling iid from the density w . Right column: the CSA method.

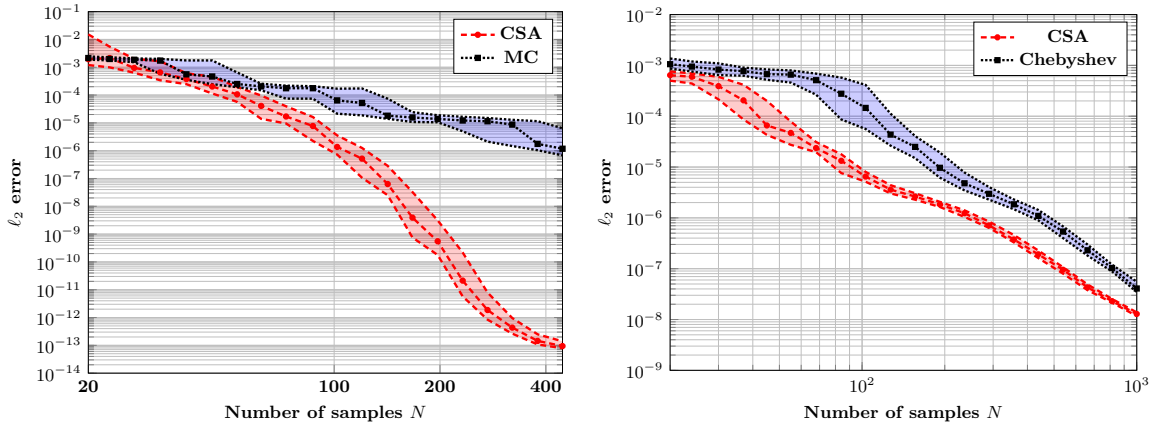


Figure 3: Results from [P9]. Compressive sampling ℓ_2 error in expansion coefficients, solving a parameterized diffusion equation with the CSA algorithm versus standard algorithms. Left: $d = 2$ with w an asymmetric Beta density. Right: $d = 20$ with w a uniform density.

3.1.3 Optimal quadrature rules: characterization and computation

Sampling is often used to approximate integrals,

$$\int_D f(x) d\mu(x) \approx \sum_{m=1}^M w_m f(x_m),$$

for nodes and weights x_m and w_m , respectively. Since function evaluations can be expensive in complex scientific models, one wishes to take as few samples as possible. This is the mathematical problem of developing optimal quadrature rules. Work in this project focused on computation of near-optimal multidimensional quadrature rules.

One theoretical accomplishment is the following characterization for polynomial quadrature rules: assume that $\Lambda \subset \mathbb{N}_0^d$ is a multi-index set so that

$$\text{span} \{\phi_j\}_{j=1}^N = \text{span} \{x^\lambda \mid \lambda \in \Lambda\} := P_\Lambda.$$

Our theory from [P11] show that any quadrature that is exact on the polynomial space, i.e.,

$$\int_D p(x) d\mu(x) = \sum_{m=1}^M w_m p(x_m), \quad p \in P_\Lambda,$$

must have M no smaller than the size of the minimal half-set of Λ , i.e.,

$$M \geq \min \left\{ |\Theta| \mid \Theta \subset \mathbb{N}_0^d, \Theta + \Theta \subseteq \Lambda \right\}.$$

Equality is achieved for optimal quadrature rules, such as univariate Gaussian quadrature rules. However, this property holds in the multivariate case, and has allowed us to construct algorithms that find near-optimal quadrature rules [P5, P11, P14] in as many as 100 dimensions.

The methodology is general, applying to non-tensor-product domains with non-standard measures μ . To illustrate this, Figure 4 shows performance of our reduced quadrature rules on a 20-dimensional function that can be well-approximated by a 2-dimensional function with non-tensorial weight and domain. Our reduced quadrature rule outperforms the accuracy of popular alternatives like sparse grids or Monte Carlo methods, and requires many orders of magnitude fewer data points. Figure 5 shows that such reduced quadrature rules are very effective for real-world topology optimization problems: We obtain an excellent approximation to the compliance PDF with approximately 20% of the effort of a more standard sparse grid procedure.

Work also focused on designing hierarchical (nested) quadrature rules using the technique of *designed quadrature* [P14]. Nested quadrature rules are useful since refinement is economical; we illustrate this in Figure 6, where the error committed by nested quadrature rules is smaller for a fixed computational budget compared to standard sparse grid constructions. Details of the approach and more examples are included in [S4].

3.2 Application: sampling in multifidelity approximation

A persistent challenge in uncertainty quantification is estimating the effect of uncertain parameters on quantities of interest. A common approach to understanding the effect of a parameter is to evaluate

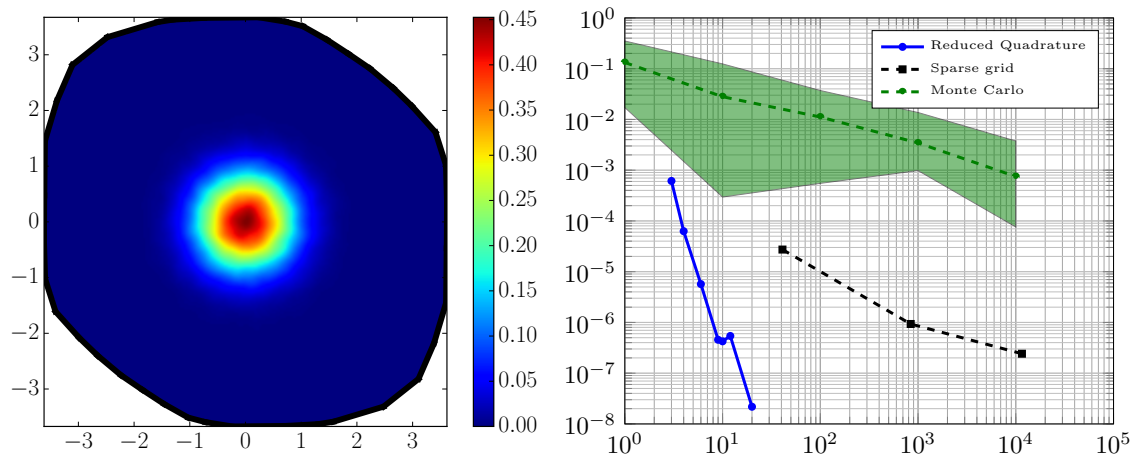


Figure 4: [P11]. Estimation of average species concentration in a chemical reaction network. (Left) The zonotope defining the domain of integration and the joint probability density of two variables for an approximate two-dimensional ridge function of 20 ambient variables. (Right) The convergence of the error in the mean value of the ridge function computed using reduced quadrature rules over the two-dimensional zonotope compared against more standard quadrature rules in the full 20-dimensional space.

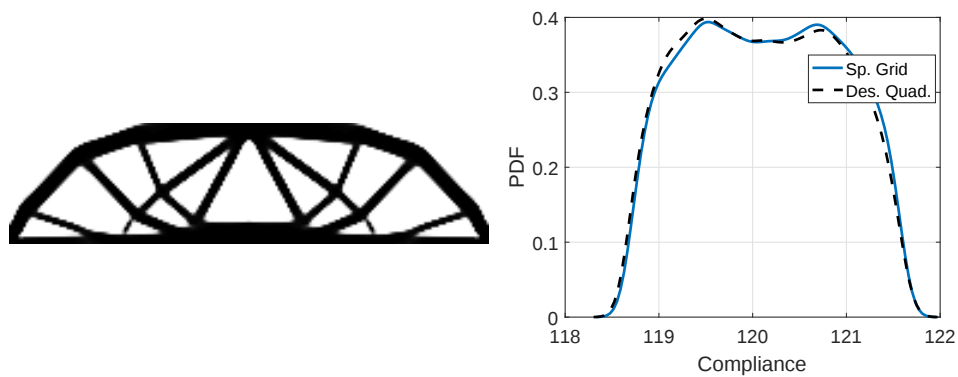


Figure 5: [P14] Robust topology design of the MBB beam (left) Compliance PDF with Sparse Grid versus near-optimal designed quadrature (right). The designed quadrature approach requires approximately 20% computational effort compared to a sparse grid approach.

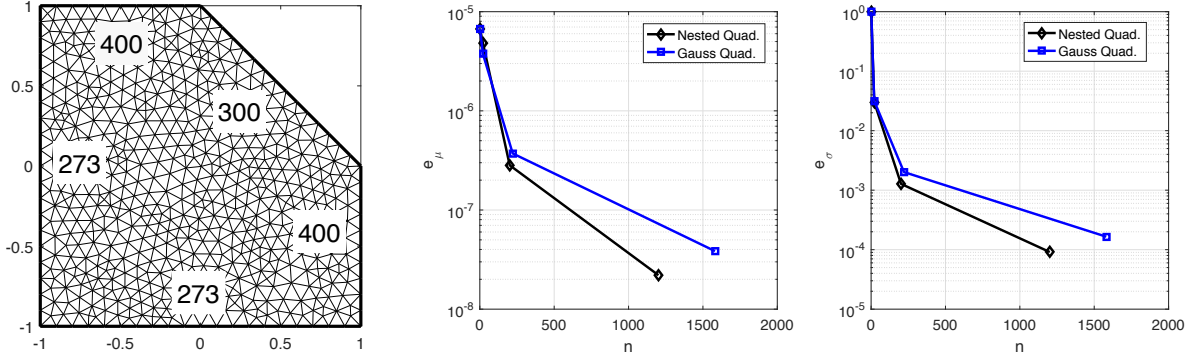


Figure 6: [S4]. A parametric elliptic partial differential equation is discretized in two dimensions on a finite element mesh (left). The parameter ($d = 10$) dictates the diffusion coefficient through a Karhunen-Loeve-type expansion. Treating the parameter as a random variable, the mean and standard deviation of the solution at a particular mesh point can be computed. The error in this computation as a function of the number of evaluations in a Gauss-type Smolyak sparse grid is compared against the error committed by the nested rules (center, right).

an ensemble of simulations for various parameter values. This approach is reasonable when multiple runs of a simulation model or experiment are easily obtained. Unfortunately, a simulation model or discretization that is more true-to-life, or has *higher fidelity*, requires additional computational resources due to, for example, increased number of discrete elements or more expensive modeling of complex phenomena. For these reasons a *high-fidelity* model simulation can incur significant computation cost, and repeating such a simulation for a sufficient number of times to understand parameter effects can quickly become infeasible.

Recent work has introduced an effective solution to this problem by using multiple fidelities of simulation models where less-expensive, lower-fidelity versions of the high-fidelity model are used to learn the parametric structure of a simulation. This structure is utilized to choose a small parameter ensemble at which high-fidelity runs are assembled, providing insight into the finer details and effects of the simulation parameters [1, 2]. One important decision in this allocation process is which parameter values should be used in the less costly low- and medium-fidelity simulations, and which values are worth the more costly high-fidelity simulation. This decision becomes extremely important for situations where it is physically impossible to run the high-fidelity simulation more than a small number of times, say $\mathcal{O}(10)$ times.

Due to the high computational cost, the number of high fidelity simulations is limited to m , with $m = \mathcal{O}(10)$ being a common bound. In contrast, the low-fidelity model is much more computationally affordable with $n \gg m$ low-fidelity simulations available.

Let $\gamma = \{z_1, \dots, z_n\}$ be a set of sample points in D , which define matrices

$$\begin{aligned} \mathbf{a}_j^L &= u^L(z_j), & \mathbf{A}^L &= [\mathbf{a}_1^L \quad \mathbf{a}_2^L \quad \dots \quad \mathbf{a}_n^L] \in \mathbb{R}^{d^L \times n}, \\ \mathbf{a}_j^H &= u^H(z_j), & \mathbf{A}^H &= [\mathbf{a}_1^H \quad \mathbf{a}_2^H \quad \dots \quad \mathbf{a}_n^H] \in \mathbb{R}^{d^H \times n}. \end{aligned}$$

Let $S \subset [n]$ denote a generic set of column indices, and for $\mathbf{A} \in \mathbb{R}^{d \times n}$ having columns $\mathbf{a}_j$ we define

$$S = \{j_1, \dots, j_{|S|}\}, \quad \mathbf{A}_S = [\mathbf{a}_{j_1} \quad \mathbf{a}_{j_2} \quad \dots \quad \mathbf{a}_{j_{|S|}}].$$

Earlier work by the PI [1] showed that the structure of $\mathbf{A}^L$ can be used to identify a small number of column indices $S \subset [n]$, $|S| = m$, so that $\mathbf{A}_S^L$ can be used to form a rank- m approximation to $\mathbf{A}^L$, and $\mathbf{A}_S^H$ can be used to form a rank- m approximation to $\mathbf{A}^H$. Precisely, they form the approximations

$$\mathbf{A}^L \approx \mathbf{A}_S^L (\mathbf{A}_S^L)^\dagger \mathbf{A}^L, \quad \mathbf{A}^H \approx \mathbf{A}_S^H (\mathbf{A}_S^L)^\dagger \mathbf{A}^L. \quad (3)$$

An important observation in the above approximations is that the representation for $\mathbf{A}^H$ requires *only* $\mathbf{A}_S^H$, i.e., it only requires $m = |S|$ evaluations of the high-fidelity model. The construction of S is performed in a greedy fashion, precisely as the first m ordered pivots in a pivoted Cholesky decomposition of $(\mathbf{A}^L)^T \mathbf{A}^L$. While $\mathbf{A}^H$ only represents u^H on a discrete set γ , the procedure above forms the approximation

$$u^H(z) \approx \sum_{i=1}^m c_i u^H(z_{j_i}), \quad (4)$$

where the coefficients c_i are the expansion coefficients of $u^L(z)$ in a least-squares approximation with the basis $\{u^L(z_{j_i})\}_{i=1}^m$. Thus, the high-fidelity simulation may actually be evaluated at any location $z \in D$ if $u^L(z)$ is known. This procedure has the following advantages:

- Once m high-fidelity simulations have been computed and stored, an approximation (4) to the high-fidelity model is constructed having the computational complexity of only the low-fidelity model, cf. (4).
- The subset S is identified via analysis of the inexpensive low-fidelity model, so that a very large set γ may be used to properly capture the parametric variation over D .
- It is not necessary for the spatiotemporal features of the low-fidelity to mimic those of the high-fidelity model. Section 6.2 in [1] reveals that u^L and u^H may actually be entirely disparate models, yet the approximation (4) can be accurate.

While the applicability of this technique is now well-understood, strong error estimates were previously unwieldy and difficult to use. Work during this project focused on making this technique more practical, providing useful computational tools for practitioners to ascertain the error committed by this approach [P8]. This resulted in development of a new, *computable* error bound that uses already-available information and is useful for practitioners. The bound can successfully predict error committed by a multifidelity technique, see Figure 7.

3.2.1 Convergence acceleration of multifidelity models

One major application was the realization that the multifidelity scheme in Section 3.2 can be used to accelerate convergence of the high-fidelity model, thereby attaining an error that is superior even to the high-fidelity model. Such a technique requires smoothness of the underlying modeling so that sequence transformation techniques (such as Richardson extrapolation) can be effective.

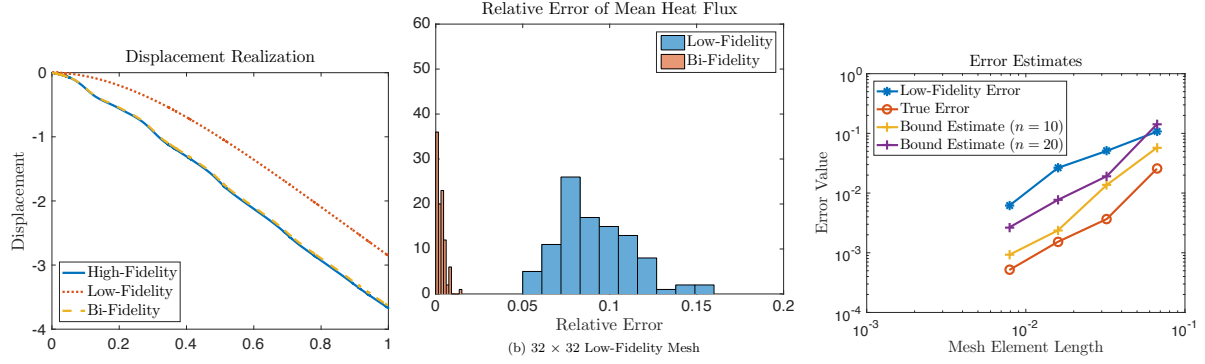


Figure 7: [P8] Left: Illustration of the effectiveness of the multifidelity approach. Center: Histogram of errors committed by the low- and multi-fidelity surrogates sampled over an ensemble of parameter values. Right: behavior of the actual error of the multifidelity surrogate relative to the high-fidelity model compared against the computable error bound.

Work during the project resulted in a rigorous theoretical error estimate demonstrating that a procedure that exploits sequence transformation is possible [P15]. Numerical results that demonstrate the possibility of such an approach were also generated and successfully achieve superior order of accuracy compared to the high-fidelity model, see Figure 8.

3.3 Application: Sampling for reduced order modeling

A lesson learned from early years of the project show that, for a given set of interpolation points $\{x_1, \dots, x_N\}$, generating new sampling points via the method

$$x_{N+1} := \operatorname{argmax}_{x \in D} \sum_{n=1}^N |\ell_n(x)|, \quad \ell_m(x_j) = \delta_{m,j}, \quad m, j = 1, \dots, N,$$

where $\{\ell_n\}_{n=1}^N$ are cardinal Lagrange interpolants. This above procedure is a greedy rule that chooses new points based on the Lebesgue function from interpolation theory. While Lagrange interpolants and Lebesgue functions are typically considered in the context of polynomial approximation, one can equally well consider this approach for non-polynomial approximations.

Later in the project, we focused on applying this technique for reduced-basis-based model order reduction of parametric PDE's. One major challenge with reduced basis methods is selection of parametric points that can be used to define a reduced basis space via solution snapshots of a parametric PDE. Work during the last period demonstrated that Lebesgue-function-based greedy sampling can be used to select a sequence of parametric points. We emphasize that this procedure utilizes *non-polynomial* approximations that are implicitly defined by the reduced order model, and in particular works well for *discontinuous* parameter dependence [P3]. Numerical results showing this are given in Figure 9. This technique is an entirely novel approach to selecting snapshots for model reduction, and we expect it to be very useful for practitioners since it is easy to implement.

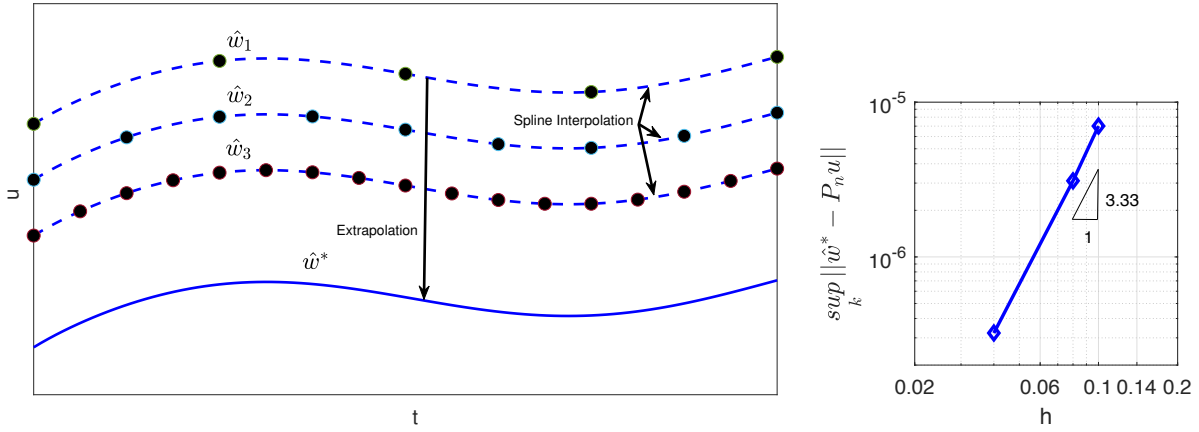


Figure 8: [P15] Multifidelity models for smooth responses can be used to effect convergence acceleration (Left), resulting in accuracy superior even to the high-fidelity model. Work during this period provided theory demonstrating that this was possible, and provided numerical results (Right) showing that, e.g., second-order methods can be accelerated to achieve third-order convergence using standard sequence transformation techniques.

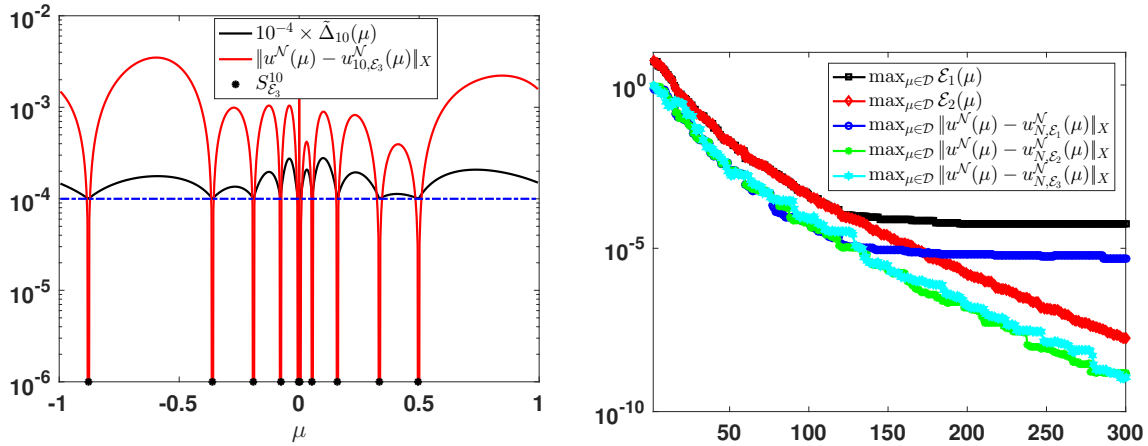


Figure 9: [P3] New, computable error estimates for general equations can help to both overcome finite-precision error stagnation and lack of more mathematically rigorous error estimates for parametric model order reduction. The new error estimates (labeled $\Delta_{10}(\mu)$) yield parametric error curves that behave similarly to more traditional error estimates, but are much easier to compute (left). The actual error committed by a model order reduction strategy using the new estimates (labeled $\mathcal{E}_3$) are comparable to those using more mathematically technical estimates (labeled $\mathcal{E}_1$ and $\mathcal{E}_2$), and can overcome finite-precision limitations that arise from the computational strategy for the mathematical estimates (right).

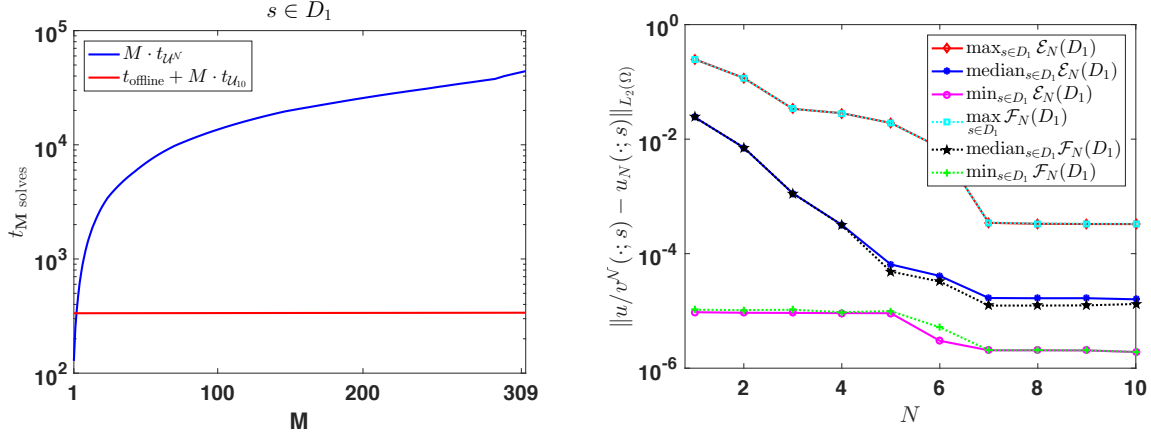


Figure 10: [S1] Model order reduction for nonlocal partial differential equations can be efficiently accomplished using the new sampling techniques of section 3.3. For the (spectral definition of the) fractional Laplace equation, the *cumulative* time required for model order reduction construction + M queries of the reduced order model is orders of magnitude smaller than M queries of the full model (left). The error committed by the reduced order model is also very accurate (right).

3.3.1 Model reduction for nonlocal equations

The novel sampling techniques from Section 3.3 were used to construct reduced order models for nonlocal differential equations. Specifically, the PDE

$$(-\Delta)^s u = f,$$

with homogeneous Dirichlet boundary conditions on $\Omega \subset \mathbb{R}^2$ was considered. Here, the order $s \in (0, 1)$ of the Laplace operator is fractional. The fact that this equation is nonlocal makes numerical discretizations expensive, and thus motivates the need for reduced order models. Using reduced basis methods, a reduced order model for the above equation can be constructed with relative ease [S1], see Figure 10.

The procedure for constructing reduced order models for nonlocal equations is one of the first attempts at creating low-rank surrogates for solutions to nonlocal equations.

3.4 Application: Topology optimization

Many structural design engineering problems can be attacked via a topology optimization framework, wherein the geometry is generated to optimize a structural metric (such as stiffness) subject to volume and/or mass constraints. Such a problem is extremely challenging due to the high-dimensional space of the design parameter. In this project we used topology optimization as a capstone application for many of our methods, in particular the adaptive sampling strategies. We summarize below the advances made in the publications [P16, P20, P14], [S5].

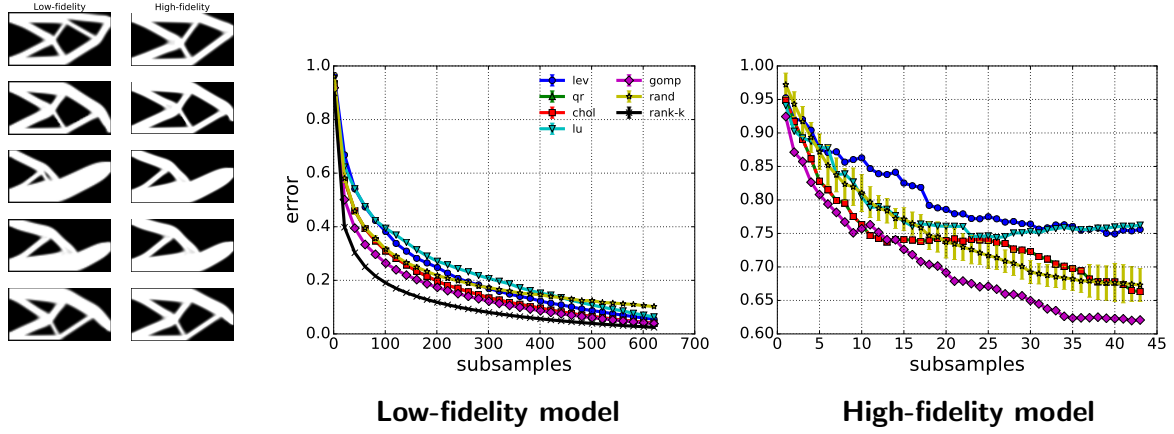


Figure 11: [P20]. **Structure topology optimization:** Reconstruction error among subset methods. *Left:* Realizations of the low-fidelity and high-fidelity topology optimization solver. *Middle:* the reconstruction error of the rest of the low-fidelity simulation dataset when using the indicated subset chosen using various methods. *Right:* the reconstruction error for the high-fidelity simulations with the subset chosen using the low-fidelity samples. The GOMP procedure outperforms methods except the best rank- k approximation, which is not computable in practical situations.

3.4.1 Machine learning samplers

In the context of the multi-fidelity procedure, certain machine learning optimization algorithms applied to the problem of adaptive sampling results in higher accuracy compared to many other sampling strategies. We show in [P20] that a type of group orthogonal matching pursuit (GOMP) algorithm for feature selection creates a more accurate multifidelity approximation – the GOMP procedure even outperforms a convolutional neural network. In Figure 11 we show results from [P20] showing performance of the GOMP algorithm on a nontrivial topology optimization problem.

3.4.2 Design under uncertainty

In [P16] we demonstrated that the ideas of adaptive sampling via multi-fidelity methods are extensible to design under uncertainty paradigms. In these problems, we seek to solve

$$\min \mu(\rho) + \lambda \sigma(\rho) \quad \text{subject to} \quad V(\rho) \leq V_0,$$

where μ and σ represent the mean and standard deviation of a structural performance metric, which is random due to fluctuations in the design or construction process. The design parameters ρ represent material mass or volume fractions over the ambient geometry, and V represents a volume or mass of the design ρ . The constants λ and V_0 are problem-dependent scalars that define the objective and constraints. In Figure 12 we demonstrate that low-fidelity models can be used to construct an adaptive set of samples in parameter space from which a multi-fidelity surrogate can be constructed and used to create designs. In Figure 13 we demonstrate this procedure on a three-dimensional design under uncertainty problem, showing both its speed and accuracy.

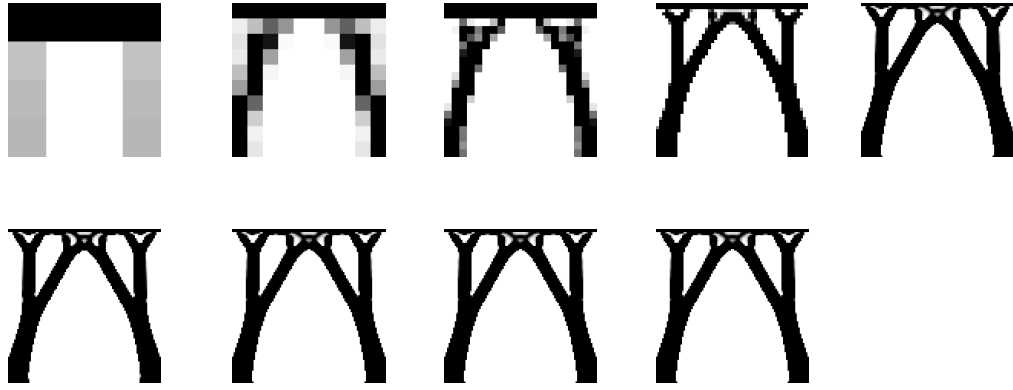


Figure 12: [P16]. Topology Optimization and Design under uncertainty using adaptive sampling from this project. Top: single-fidelity optimization with 4×4 , 10×10 , 20×20 , 50×50 and 100×100 meshes (top row). Bottom row: sampling-based bi-fidelity optimization with low-fidelity meshes of size 4×4 , 10×10 , 20×20 , 50×50 meshes (bottom row). The top right simulation requires approximately 15 times the computational cost compared to the lower left simulation.

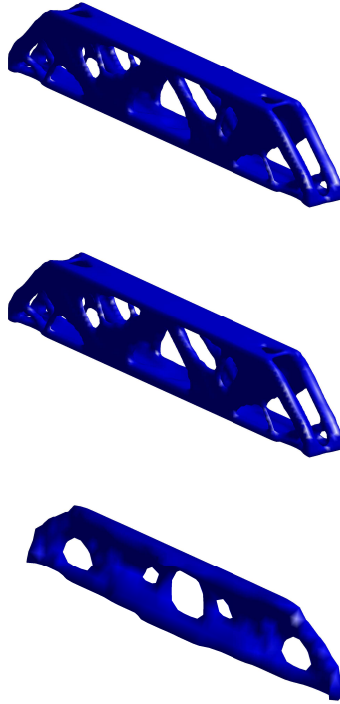


Figure 13: [P16]. Topology Optimization and Design under uncertainty using adaptive sampling from this project. Top: high-fidelity optimization. Middle: adaptive multi-fidelity optimization. Bottom row: low-fidelity optimization. The top simulation requires approximately 1500 times more costly than the middle simulation, despite the lack of appreciable difference between the two.

References

- [1] Akil Narayan, Claude Gittelsohn, and Dongbin Xiu. A Stochastic Collocation Algorithm with Multifidelity Models. *SIAM Journal on Scientific Computing*, 36(2):A495–A521, 2014.
- [2] Xueyu Zhu, Akil Narayan, and Dongbin Xiu. Computational Aspects of Stochastic Collocation with Multifidelity Models. *SIAM/ASA Journal on Uncertainty Quantification*, 2(1):444–463, 2014.