

Quadrature Sampling of Parametric Models with Bi-fidelity Boosting

Nuojin Cheng^{*1}, Osman Asif Malik^{*2}, Yiming Xu^{*3}, Stephen Becker¹, Alireza Doostan⁴,
and Akil Narayan⁵

¹Department of Applied Mathematics, University of Colorado Boulder (nuojin.cheng@colorado.edu,
stephen.becker@colorado.edu)

²Applied Mathematics & Computational Research Division, Lawrence Berkeley National Laboratory (oamalik@lbl.gov)

³Corporate Model Risk, Wells Fargo (yiming.xu@wellsfargo.com)

⁴Smead Aerospace Engineering Sciences Department, University of Colorado Boulder (alireza.doostan@colorado.edu)

⁵Scientific Computing and Imaging Institute, and Department of Mathematics, University of Utah (akil@sci.utah.edu)

Abstract

Least squares regression is a ubiquitous tool for building emulators (a.k.a. surrogate models) of problems across science and engineering for purposes such as design space exploration and uncertainty quantification. When the regression data are generated using an experimental design process (e.g., a quadrature grid) involving computationally expensive models, or when the data size is large, sketching techniques have shown promise to reduce the cost of the construction of the regression model while ensuring accuracy comparable to that of the full data. However, random sketching strategies, such as those based on leverage scores, lead to regression errors that are random and may exhibit large variability. To mitigate this issue, we present a novel boosting approach that leverages cheaper, lower-fidelity data of the problem at hand to identify the *best* sketch among a set of candidate sketches. This in turn specifies the sketch of the intended high-fidelity model and the associated data. We provide theoretical analyses of this bi-fidelity boosting (BFB) approach and discuss the conditions the low- and high-fidelity data must satisfy for a successful boosting. In doing so, we derive a bound on the residual norm of the BFB sketched solution relating it to its ideal, but computationally expensive, high-fidelity boosted counterpart. Empirical results on both manufactured and PDE data corroborate the theoretical analyses and illustrate the efficacy of the BFB solution in reducing the regression error, as compared to the non-boosted solution.

1 Introduction

Computational models are becoming central tools in analysis, design, and prediction. In these models, input parameters are often modeled as a random vector $\mathbf{p}$ to account for either uncertainty in precise values of these parameters, or as a means to model variability of parameters in order to assess robustness of an output [LMK10; Smi13]. We consider such types of models given a (possibly non-linear) parameter-to-output map,

$$\mathbf{b} = \mathcal{T}(\mathbf{p}), \quad \mathcal{T} : \mathbb{R}^q \rightarrow \mathbb{R}.$$

A canonical example is when $\mathcal{T}$ is a measurement functional (e.g., the spatial average) operating on the solution to an elliptic partial differential equation (PDE) whose formulation contains random variables that, e.g., parameterize the diffusion coefficient. Hence, $\mathcal{T}$ is the composition of a measurement functional with

^{*}Equal contribution.

the solution map of a parametric PDE. By placing a probability distribution on $\mathbf{p}$ that reflects a model of uncertainty, the goal of forward uncertainty quantification (UQ) is to quantify the resulting randomness in $b(\mathbf{p})$, frequently via statistics. Since explicit formulas revealing the dependence of b on $\mathbf{p}$ are typically not available, one resorts to approximations. One such sampling-based approach that we focus on is that of polynomial chaos (PC) methods [GS03; XK02] using variants of stochastic collocation [XH05].

In this paper we consider building emulators for forward UQ via a non-intrusive least squares-based PC strategy. More precisely, we assume an *a priori* form for an emulator b_V :

$$b(\mathbf{p}) \approx b_V(\mathbf{p}) := \sum_{j=1}^d x_j^* \psi_j(\mathbf{p}), \quad V := \text{span}\{\psi_1, \dots, \psi_d\}, \quad (1.1)$$

where ψ_j are fixed, known functions (in PC approaches they are multivariate polynomial functions of $\mathbf{p}$), and the coefficients x_j^* must be determined. We identify these coefficients through data collected from evaluating b on a prescribed quadrature rule $\{(\mathbf{p}_n, w_n)\}_{n=1}^N$, with quadrature nodes $\mathbf{p}_n$ and positive weights w_n . The coefficients x_j^* are then chosen as the solution to a quadrature-based least squares problem,

$$\mathbf{x}^* = \arg \min_{\mathbf{x} \in \mathbb{R}^d} \|\mathbf{A}\mathbf{x} - \mathbf{b}\|_2^2, \quad \mathbf{A}(n, j) = \sqrt{w_n} \psi_j(\mathbf{p}_n), \quad \mathbf{b}(n) = \sqrt{w_n} b(\mathbf{p}_n), \quad (1.2)$$

where $\mathbf{A} \in \mathbb{R}^{N \times d}$ is referred to as the *design matrix* of the problem. Once $\mathbf{x}^*$ is computed, the emulator b_V is easily manipulated and computationally analyzed to compute (approximate) statistics for b or the sensitivity of b to each entry of $\mathbf{p}$. The challenge with this approach is that when $\dim \mathbf{p} = q \gg 1$, then designing an appropriately accurate quadrature rule requires $N \gg 1$ samples of b , which is prohibitively expensive when such evaluations amount to PDE solutions. (For example a q -dimensional tensorized Gaussian quadrature rule with n points per dimension requires $N = n^q$ points.)

In this paper, we describe one strategy to mitigate this cost via a procedure that combines statistical boosting ideas from theoretical computer science (see, e.g., [Mah11, Sec. 7.2] and [Woo14, Sec. 2.3]) with bi-fidelity strategies in UQ. More precisely, our approach boosts on the randomness of a *sketching* operator $\mathbf{S} \in \mathbb{R}^{m \times N}$ that is used to approximately solve (1.2):

$$\hat{\mathbf{x}} = \arg \min_{\mathbf{x} \in \mathbb{R}^d} \|\mathbf{S}\mathbf{A}\mathbf{x} - \mathbf{S}\mathbf{b}\|_2^2.$$

Without *a priori* knowledge of $\mathbf{b}$, a deterministic sketch with $m < N$ generally is not robust to adversarial vectors $\mathbf{b}$ that result in a large residual for $\hat{\mathbf{x}}$ relative to the residual for $\mathbf{x}^*$. However, in general scenarios one can identify constructive *probabilistic* models for $\mathbf{S}$ where sketches of near-optimal size, $m \gtrsim d \log d / (\epsilon \delta)$, ensure

$$\|\mathbf{A}\hat{\mathbf{x}} - \mathbf{b}\|^2 \leq (1 + \epsilon) \|\mathbf{A}\mathbf{x}^* - \mathbf{b}\|^2 \text{ with probability } \geq 1 - \delta.$$

We provide a more detailed discussion of existing sketching guarantees in section 2.2, in particular for *row sketches*, for which computing $\mathbf{S}\mathbf{b}$ requires knowledge of only m entries of $\mathbf{b}$, rather than all N entries. While random sketching provides attractive guarantees when $m \ll N$, it is still random and hence is subject to randomness in performance, and “failure” events can occur with nonzero probability δ . Naive statistical boosting mitigates this issue by generating several (say L) sketches and choosing the one that yields the smallest residual. However, this requires generating Lm entries of $\mathbf{b}$, which can be computationally expensive when each evaluation is an expensive PDE solve. Our approach attacks this problem in the sketch selection boosting phase by replacing $\mathbf{b}$ with an approximate, low-fidelity version from which collecting Lm samples is computationally feasible. Once a “good” sketch is identified in the boosting phase, we solve the sketched least squares problem using the corresponding sketch of the original data $\mathbf{b}$.

Thus, we assume availability of and leverage a *low-fidelity* model $\tilde{b}(\mathbf{p})$. For example, $\tilde{b}$ may correspond to using a discretized PDE solver with a mesh coarser than the one which produces accurate realizations of b , or to model approximations such as Reynolds-averaged Navier Stokes solvers, or to solutions computed with

arithmetic in lower precision compared to samples for b . Although $\tilde{b}$ may be untrusted as a replacement for b , it can be used to extract some useful information about b , as is done in by-now standard multi-fidelity approaches [PWG18]. Throughout this paper, we assume the bi-fidelity setup, i.e., two levels of fidelity, and also that the cost of evaluating $\tilde{b}$ is much less than the corresponding cost for b ; both of these are common practical assumptions [DGRH07; NGX14; ZNX14; Fai+20; New+22].

1.1 Contributions of this article

The contributions of this article are as follows:

- We propose a new bi-fidelity boosting (BFB) algorithm to compute an approximation to $\mathbf{x}^*$. The procedure, given in Algorithm 2, computes the solution of a *sketched* least squares problem, where the sketch matrix is identified by a boosting procedure on a low-fidelity data vector $\tilde{\mathbf{b}}$. The sketching approach reduces the required sample complexity from N evaluations of b to $\sim d \log d$ samples of b , which can be a significant saving. The boosting procedure requires $\sim L d \log d$ evaluations of the low-fidelity model $\tilde{b}$, where, in the language of statistical learning, L is the number of weak learners used in the boosting procedure. When $\tilde{b}$ costs substantially less than b , this cost for collecting the boosting data is negligible.
- We provide a theoretical analysis of BFB under certain assumptions, which provides quantitative bounds on the residual of the BFB solution $\hat{\mathbf{x}}_{\text{BFB}}$ relative to the full, computationally expensive solution $\mathbf{x}^*$ (see Theorems 3.2 and 3.4). We also provide some asymptotic bounds on the correlation between the low- and high-fidelity solutions in a certain sense (see Theorem 3.5). Finally, we provide concrete computational strategies to ensure that the required assumptions of BFB hold (see Theorem 3.11).
- We investigate the numerical performance of BFB when combined with several different sampling strategies and compare the performance to the corresponding sampling strategies without boosting. We also demonstrate using real-world problems that the assumptions required for BFB’s theoretical analysis frequently hold in practice.

The idea of sketching for least squares solutions has a substantial history in the computer science and numerical linear algebra communities [Mah11; Woo14]. Our use of sparse row sketches of size $\sim d$ is identical to existing methods for leverage score-based [Mah11], Gaussian-sketch based [MT20], and volume-maximizing sketching [DW18; DWH18]. In addition, boosting for least squares problems is also not a new idea [HNP22]. However our combination of these approaches in a bi-fidelity setting is new to our knowledge, and our analysis in this bi-fidelity context provides novel, non-trivial insight into the algorithm performance.

The rest of this manuscript is organized as follows. Section 2 introduces the notation we use and provides some background material on various sketching approaches in least squares approximation. Section 3 presents the BFB algorithm along with its theoretical analysis. Section 4 contains numerical experiments which illustrate various aspects of the BFB approach. We conclude the present study in Section 5. The paper also contains several appendices. Appendix A provides a brief introduction to the sampling approach that we proposed in [Mal+22] and which we make use of in this paper. Appendices B and C contain some proofs that have been left out of the main text.

2 Preliminaries

For the interest of clarity and completeness, we next introduce the notation used throughout the manuscript and introduce four sampling strategies to sketch the least squares problem (1.2), namely, sampling via column-pivoted QR, leverage scores, volume maximization, and Gaussian distribution.

2.1 Notation

Matrices are denoted by bold upper-case letters (e.g., $\mathbf{A}$), vectors are denoted by bold lower-case letters (e.g., $\mathbf{x}$) and scalars by lower case regular and Greek letters (e.g., a and α). Entries of matrices and vectors are indicated in parentheses. For example, $\mathbf{A}(i, j)$ is the entry on position (i, j) in $\mathbf{A}$ and $\mathbf{a}(i)$ is the i th entry in $\mathbf{a}$. A colon is used to denote all entries along a mode of a matrix. For example, $\mathbf{A}(i, :)$ is the i th row of $\mathbf{A}$ represented as a row vector. For a set of indices $\mathcal{J}$, $\mathbf{A}(\mathcal{J}, :)$ denotes the submatrix $(\mathbf{A}(j, :))_{j \in \mathcal{J}}$ and $\mathbf{a}(\mathcal{J})$ denotes the subvector $(\mathbf{a}(j))_{j \in \mathcal{J}}$.

The *compact SVD* of a matrix $\mathbf{A}$ takes the form $\mathbf{A} = \mathbf{U}\mathbf{\Sigma}\mathbf{V}^\top$, where $\mathbf{U}$ and $\mathbf{V}$ have $\text{rank}(\mathbf{A})$ columns and $\mathbf{\Sigma}$ is of size $\text{rank}(\mathbf{A}) \times \text{rank}(\mathbf{A})$. The *pseudoinverse* of $\mathbf{A}$ is denoted by $\mathbf{A}^\dagger \stackrel{\text{def}}{=} \mathbf{V}\mathbf{\Sigma}^{-1}\mathbf{U}^\top$. For a matrix $\mathbf{U}$ with orthonormal columns, we use $\mathbf{U}_\perp$ to denote an orthonormal complement of $\mathbf{U}$, i.e., $\mathbf{U}_\perp$ is any matrix such that $[\mathbf{U}, \mathbf{U}_\perp]$ is square and has orthonormal columns. We use $\mathbf{P}_\mathbf{A} \stackrel{\text{def}}{=} \mathbf{A}\mathbf{A}^\dagger = \mathbf{U}\mathbf{U}^\top$ to denote the *orthogonal projection* onto $\text{range}(\mathbf{A})$, where $\mathbf{U} = \text{orth}(\mathbf{A})$ is a(ny) matrix whose columns are an orthonormal basis for $\text{range}(\mathbf{A})$, e.g., via the compact SVD or QR decomposition of $\mathbf{A}$. The determinant of $\mathbf{A}$ is denoted by $\det(\mathbf{A})$. For a positive integer n , we use the notation $[n] \stackrel{\text{def}}{=} \{1, 2, \dots, n\}$. We use $\mathbf{a}_\mathcal{P}$ to denote a vector $\mathbf{a} \neq \mathbf{0}$ rescaled to unit length:

$$\mathbf{a}_\mathcal{P} = \frac{\mathbf{a}}{\|\mathbf{a}\|_2}.$$

We also introduce two notions of correlation: for given deterministic vectors $\mathbf{a}, \mathbf{b} \neq \mathbf{0}$, we define the correlation between them as the cosine of the angle separating them:

$$\text{corr}(\mathbf{a}, \mathbf{b}) \stackrel{\text{def}}{=} \frac{\langle \mathbf{a}, \mathbf{b} \rangle}{\|\mathbf{a}\|_2 \|\mathbf{b}\|_2},$$

where $\langle \cdot, \cdot \rangle$ denotes the Euclidean inner product. We will also require Pearson's correlation coefficient, which is widely used in statistics. For two (non-constant) random variables X and Y with bounded second moments defined on the same probability space, their correlation is defined as

$$\text{corr}(X, Y) \stackrel{\text{def}}{=} \frac{\mathbb{E}[(X - \mathbb{E}[X])(Y - \mathbb{E}[Y])]}{\sqrt{\mathbb{V}[X]\mathbb{V}[Y]}},$$

where $\mathbb{E}[\cdot]$ and $\mathbb{V}[\cdot]$ are, respectively, the mathematical expectation and variance operators. Note that our notation $\text{corr}(\cdot, \cdot)$ is overloaded, operating differently on vectors and (random) scalars. The context of use in what follows should make it clear which definition above is used.

We will use the following notation to denote the minimum of the least squares objective in (1.2):

$$r(\mathbf{A}, \mathbf{b}) \stackrel{\text{def}}{=} \min_{\mathbf{x}} \|\mathbf{A}\mathbf{x} - \mathbf{b}\|_2 = \|\mathbf{A}\mathbf{x}^* - \mathbf{b}\|_2, \quad (2.1)$$

where $\mathbf{x}^*$ is defined as in (1.2).

2.2 Sketching of least squares problems

Solving the problem (1.2) using standard methods (e.g., via the QR decomposition) costs $\mathcal{O}(Nd^2)$ ¹. When N is large, this may be prohibitively expensive. A popular approach to address this issue is to apply a *sketch operator* $\mathbf{S} \in \mathbb{R}^{m \times N}$ where $m \ll N$ to both $\mathbf{A}$ and $\mathbf{b}$ in (1.2) in order to reduce the size of the problem:

$$\hat{\mathbf{x}} \stackrel{\text{def}}{=} \arg \min_{\mathbf{x} \in \mathbb{R}^d} \|\mathbf{S}\mathbf{A}\mathbf{x} - \mathbf{S}\mathbf{b}\|_2. \quad (2.2)$$

This approach has two benefits: (i) If $\mathbf{S}$ is a row-sketch, i.e., has only a small number of non-zero columns, then $\mathbf{S}\mathbf{b}$ requires knowledge of only a small number of entries of $\mathbf{b}$, and (ii) the cost of solving this smaller

¹In our context, we have $N > d$; see Assumption 3.1.

problem is $\mathcal{O}(md^2)$, a substantial reduction from $\mathcal{O}(Nd^2)$ when $m \ll N$. Analogously to (2.1), we will use the following to denote the least squares objective value for the approximate solution:

$$r_{\mathbf{S}}(\mathbf{A}, \mathbf{b}) \stackrel{\text{def}}{=} \|\mathbf{A}\hat{\mathbf{x}} - \mathbf{b}\|_2. \quad (2.3)$$

The goal is for the approximation $\hat{\mathbf{x}}$ to yield a residual “close” to the optimal residual of the full problem (1.2),

$$r(\mathbf{A}, \mathbf{b}) \approx r_{\mathbf{S}}(\mathbf{A}, \mathbf{b}),$$

which is typically achieved if m is “large enough”. The following definition makes this more precise.

Definition 2.1 ((ε, δ) pair condition). Let $\mathbf{S} \in \mathbb{R}^{m \times N}$ be a random matrix. Given $\mathbf{A} \in \mathbb{R}^{N \times d}$, $\mathbf{b} \in \mathbb{R}^N$, and $\varepsilon, \delta > 0$, the distribution of $\mathbf{S}$ is said to satisfy an (ε, δ) pair condition for $(\mathbf{A}, \mathbf{b})$ if, with probability at least $1 - \delta$, both conditions,

$$\text{rank}(\mathbf{S}\mathbf{A}) = \text{rank}(\mathbf{A}) \quad \text{and} \quad r_{\mathbf{S}}(\mathbf{A}, \mathbf{b}) \leq (1 + \varepsilon) r(\mathbf{A}, \mathbf{b}), \quad (2.4)$$

hold simultaneously, where $r(\mathbf{A}, \mathbf{b})$ and $r_{\mathbf{S}}(\mathbf{A}, \mathbf{b})$ are defined as in (2.1) and (2.3), respectively.

Note that one can only ask for the above condition with probability less than 1: For any sketch with $m < N$, there are vectors $\mathbf{b}$ for which the residual bound condition in (2.4) can be violated. Such a condition can be satisfied with $m < N$ samples; see sections 2.2.2, 2.2.3, and 2.2.4. Sketching operators $\mathbf{S}$ that sample a subset of the rows are of particular interest in UQ since $\mathbf{S}\mathbf{b}$ in (2.2) then requires knowledge of only a subset of entries in the vector $\mathbf{b}$, meaning that fewer samples need to be collected. In this paper, we consider three different sketching operators of this type, one of which is deterministic and two of which are random. These are described in Sections 2.2.1–2.2.3. Another popular sketching operator is the Gaussian sketching operator whose entries are appropriately scaled i.i.d. normal random variables. Applying such a random matrix to $\mathbf{b}$ requires knowledge of all entries in $\mathbf{b}$. While this makes the Gaussian sketch unsuitable for use in practice for quadrature sampling, we still consider it in some of our theoretical results since it is easier to analyze than the sampling-based sketches. Furthermore, since it is known to have excellent guarantees, it provides a nice baseline. We introduce the Gaussian sketch in Section 2.2.4.

Much research has been conducted over the last two decades on randomized algorithms in numerical linear algebra, including the problem of solving least squares problems. We only cover the basics that are relevant for this paper. For a more in-depth discussion, we refer the reader to the surveys in [HMT11; Mah11; Woo14; MT20] and the references therein.

2.2.1 Sampling via column-pivoted QR decomposition

Let $\mathbf{A}^\top \mathbf{P} = \mathbf{A}(\mathcal{J}, :)^{\top} = \mathbf{Q}\mathbf{R}$ be a column-pivoted QR (CPQR) decomposition where $\mathcal{J}$ is a length- N permutation vector. A simple deterministic heuristic for sampling m rows from $\mathbf{A}$ is to simply choose those rows corresponding to the first m entries in $\mathcal{J}$, i.e., $\mathbf{A}(\mathcal{J}(1:m), :)$. This corresponds to applying a sketch $\mathbf{S} = (\mathbf{P}(:, 1:m))^{\top}$ to $\mathbf{A}$. Such an approach has been used to sub-sample points from either tensor product quadratures [SNM17] or from random samples (approximate D-optimal design) [HD18a; DDH18; Guo+18] in the context of least squares polynomial approximation.

Recall that $\mathbf{A}$ is an $N \times d$ tall-and-skinny matrix. When $m \leq d$, the subsample is straightforward and just takes the first m entries in $\mathcal{J}$ since the list $\mathcal{J}$ contains the entries in decreasing order of importance (as approximated by the column-pivoting algorithm). When $m > d$, the situation is more subtle since the remaining entries $\mathcal{J}(d+1 : N)$ have no particular meaning and will not be useful in our row-sampling procedure. To get around this, we use the heuristic in Algorithm 1 in order to sample $m > d$ rows. The heuristic chooses the first d rows indices to be the entries in $\mathcal{J}(1:d)$ where $\mathcal{J}$ comes from the column-pivoted QR decomposition of $\mathbf{A}^\top$. The rows with indices in $\mathcal{J}(1:d)$ are then removed from $\mathbf{A}$. Another column-pivoted QR decomposition is then computed for the updated $\mathbf{A}^\top$, and the next set of d rows is chosen to be the rows of $\mathbf{A}$ corresponding to the top- d entries in the new permutation vector $\mathcal{J}$. Once again, the chosen

Algorithm 1: Heuristic for sampling via column-pivoted QR decomposition

Input: $\mathbf{A}$: design matrix; m : desired number of row samples

Output: $\mathbf{A}_s$: matrix containing m rows of $\mathbf{A}$

- 1: Initialize $\mathbf{A}_s$ to an empty matrix: $\mathbf{A}_s = []$
 - 2: **while** $m > 0$ **do**
 - 3: Compute column-pivoted QR of $\mathbf{A}^\top$: $\mathbf{A}(\mathcal{J}, :)^{\top} = \mathbf{Q}\mathbf{R}$
 - 4: Let $k = \min(d, m)$
 - 5: Append top- k rows from $\mathbf{A}$ to $\mathbf{A}_s$: $\mathbf{A}_s = [\mathbf{A}_s; \mathbf{A}(\mathcal{J}(1:k), :)]$
 - 6: Remove top- k rows from $\mathbf{A}$: $\mathbf{A} = \mathbf{A}(\mathcal{J}(k+1:\text{end}), :)$
 - 7: $m = m - k$
 - 8: **end while**
 - 9: **return** $\mathbf{A}_s$
-

rows are removed from $\mathbf{A}$. This procedure is repeated until m rows have been chosen. It is straightforward to formulate a sampling matrix $\mathbf{S}$ such that $\mathbf{SA} = \mathbf{A}_s$, where $\mathbf{A}_s$ is the output of Algorithm 1.

Since the approach in Algorithm 1 is deterministic, it cannot satisfy guarantees of the form in Definition 2.1. However, for the case $m = d$ it is possible to prove bounds on the condition number of $\mathbf{A}(\mathcal{J}(1:d), :)$; see Lemma 2.1 in [SNM17] for details.

2.2.2 Leverage score sampling

Let $\mathbf{A} = \mathbf{U}\mathbf{\Sigma}\mathbf{V}^\top$ be a compact SVD. The *leverage scores* of $\mathbf{A}$ are defined as

$$\ell_i(\mathbf{A}) \stackrel{\text{def}}{=} \|\mathbf{U}(i, :)\|_2^2 \quad \text{for } i \in [N]. \quad (2.5)$$

They take values in the range $\ell_i(\mathbf{A}) \in [d/N, 1]$ and indicate how important each row of $\mathbf{A}$ is in a certain sense. The matrix $\mathbf{U}$ can be replaced with any matrix whose columns form an orthonormal basis for $\text{range}(\mathbf{A})$ without impacting the definition in (2.5) [Woo14, Sec. 2.4]. The *coherence* of $\mathbf{A}$ is defined as

$$\gamma(\mathbf{A}) \stackrel{\text{def}}{=} \max_{i \in [N]} \ell_i(\mathbf{A}).$$

It takes values in the range $\gamma(\mathbf{A}) \in [d/N, 1]$; it is maximal when one of the leverage scores is 1 and minimal when all leverage scores are equal to d/N . Let $r \stackrel{\text{def}}{=} \sum_i \ell_i(\mathbf{A})$. The leverage score sampling distribution of $\mathbf{A}$ is defined as

$$p_i(\mathbf{A}) \stackrel{\text{def}}{=} \frac{\ell_i(\mathbf{A})}{r} \quad \text{for } i \in [N],$$

which is indeed a probability distribution as $\ell_i(\mathbf{A}) > 0$. Let $f : [m] \rightarrow [N]$ be a random map such that each $f(j)$ is independent and $\mathbb{P}\{f(j) = i\} = p_i(\mathbf{A})$ for each $j \in [m]$. The leverage score sampling sketch $\mathbf{S} \in \mathbb{R}^{m \times N}$ is defined elementwise via

$$S_{ji} = \frac{\text{Ind}\{f(j) = i\}}{\sqrt{mp_{f(j)}(\mathbf{A})}} \quad \text{for } (j, i) \in [m] \times [N], \quad (2.6)$$

where $\text{Ind}\{A\}$ is the indicator function which is 1 if the random event A occurs and zero otherwise. Algorithms and theory for leverage score sampling have been developed in a number of papers; see e.g., [DMM06; DMM08; Dri+11; Mah11; LK20] and references therein. The distribution for the leverage score sketch in (2.6) satisfies an (ε, δ) condition for $(\mathbf{A}, \mathbf{b})$ if

$$m \gtrsim d \log(d/\delta) + d/(\varepsilon\delta); \quad (2.7)$$

see Theorem 3.11 for a more detailed and slightly stronger statement.

Choosing $p_i(\mathbf{A}) = 1/N$ results in *uniform sampling*. For general matrices, there are no useful guarantees when sampling uniformly in this fashion. However, if $\mathbf{A}$ has low coherence, then uniform sampling will be close to the leverage score sampling distribution and guarantees similar to those for leverage score sampling hold. More precisely, if $\ell_i(\mathbf{A}) \leq Cd/N$ for some constant $C \geq 1$, then uniform sampling satisfies an (ε, δ) condition for $(\mathbf{A}, \mathbf{b})$ if m is chosen as in (2.7) (this is a direct consequence of, e.g., Theorem 6 in [LK20]). Notice that the difference from sampling according to the exact leverage scores is that there now is an additional constant C hidden in the lower bound on m .

In addition to a parsimonious sampling of $\mathbf{b}$, the computational complexity of the sketched least squares approach in (2.2) is a consideration. Direct sampling of the leverage score distribution via the formula (2.5) requires a matrix decomposition (e.g., QR or SVD), which costs $\mathcal{O}(Nd^2)$ effort, the same effort required to solve the original least squares problem. Drineas et al. [Dri+12] propose a procedure for computing leverage score estimates with cost $\mathcal{O}(Nd \log N)$ for any matrix $\mathbf{A}$. When $\mathbf{A}$ has particular structure it is possible to improve this considerably. Malik et al. [Mal+22] propose such a method for the case when the multivariate basis functions ψ_j in (1.1) are certain products of one-dimensional functions, which corresponds to impose certain structure on the subspace V . In the polynomial approximation setting, those structural conditions are satisfied by a large family of subspaces, including the popular tensor product, total degree, and hyperbolic cross spaces. For example, if the multivariate basis polynomials for q -dimensional inputs correspond to polynomials of at most degree k in each dimension and use n grid points per dimension (in which case $\mathbf{A}$ has $N = n^q$ rows), then the total cost of our method is at most $\mathcal{O}(qnk^2 + mq)$ for drawing m samples. This sampling approach is an ingredient in our method, so we describe the key aspects of how this sampling approach works in Appendix A and refer the reader to [Mal+22] for a more comprehensive treatment.

2.2.3 Leveraged volume sampling

Volume sampling is a technique that samples a set $\mathcal{J} \subset [N]$ of m row indices of $\mathbf{A}$ with probability proportional to the squared volume of the parallelepiped spanned by the columns of the submatrix $\mathbf{A}(\mathcal{J}, :)$, i.e., $\mathbb{P}(\mathcal{J}) \propto \det(\mathbf{A}(\mathcal{J}, :)^{\top} \mathbf{A}(\mathcal{J}, :))$. This means that, unlike for leverage score sampling, the rows are not sampled independently. This has several benefits, including that the sketched least square solution $\mathbf{A}(\mathcal{J}, :)^{\dagger} \mathbf{b}(\mathcal{J})$ is correct in expectation [DW17, Prop. 7]: $\mathbb{E}[\mathbf{A}(\mathcal{J}, :)^{\dagger} \mathbf{b}(\mathcal{J})] = \mathbf{A}^{\dagger} \mathbf{b}$. Leverage score sampling, by contrast, may produce a biased estimate of the solution vector. Despite the apparent issue of sampling from a combinatorial number of subsets of $[N]$, there are algorithms for volume sampling that run in polynomial time. Dereziński and Warmuth [DW18] propose two such algorithms, RegVol and FastRegVol. RegVol runs in $\mathcal{O}((N - m + d)Nd)$ time, and FastRegVol runs in $\mathcal{O}((N + \log(N/d) \log(1/\delta))d^2)$ time with probability at least $1 - \delta$. The dependence on N can be prohibitive in quadrature sampling since the number of (tensor-product) quadrature points N is exponential in the number of variables.

Dereziński, Warmuth, and Hsu [DWH18] propose *leveraged* volume sampling which improves on standard volume sampling in several ways. Importantly, it still retains the correctness in expectation but allows for more efficient sampling. In particular, the cost of sampling does not depend on N . Unlike standard volume sampling, the sketch distribution satisfies an (ε, δ) condition for $(\mathbf{A}, \mathbf{y})$ if $m \gtrsim d \log(d/\delta) + d/(\varepsilon\delta)$, which is on par with what leverage score sampling requires for such guarantees. Leveraged volume sampling has two stages. In the first stage, $\mathcal{O}(d^2)$ rows are chosen from $\mathbf{A}$ using a combination of leverage score sampling and rejection sampling. After that, the $\mathcal{O}(d^2)$ subset is further reduced to $\mathcal{O}(d \log(d/\delta) + d/(\varepsilon\delta))$ via standard volume sampling. In the experiments, we use FastRegVol from [DW18] for the second step. When FastRegVol is used, the cost of leveraged volume sampling is $\mathcal{O}(((d^2 + m)d^2 + mC_{\text{samp}}) \log(1/\delta))$, where C_{samp} is the cost of drawing one row index of $\mathbf{A}$ using leverage score sampling. As discussed in Section 2.2.2, the cost C_{samp} of leverage score sampling can be reduced drastically in our setting by using the structured sampling techniques from [Mal+22].

2.2.4 Gaussian sketching operator

The Gaussian sketching operator $\mathbf{S} \in \mathbb{R}^{m \times N}$ has entries that are i.i.d. Gaussian random variables with mean zero and variance $1/m$. The Gaussian sketch satisfies an (ε, δ) condition if $m \gtrsim (d/\varepsilon) \log(d/\delta)$. These results also extend to the case when the entries of $\mathbf{S}$ are sub-Gaussian; see Theorem 3.11 for further details.

The main benefit of the Gaussian sketching operator is that it allows for simple and precise theoretical analysis of procedures that use sketching as a subroutine [MT20, Remark 8.2]. This is our motivation for considering the Gaussian sketch in this paper. Computationally, it is not efficient to use Gaussian sketching for least squares problems. The reason is that computing $\mathbf{S}\mathbf{A}$ costs $\mathcal{O}(mNd)$ which is more than the $\mathcal{O}(Nd^2)$ cost of solving the original least squares problem (recall that $m > d$). As discussed earlier, an additional issue in bi-fidelity estimation is that computing $\mathbf{S}\mathbf{b}$ requires knowledge of all elements of $\mathbf{b}$ which is prohibitively expensive when that vector contains high-fidelity data.

2.3 Bi-fidelity problems

The main goal of this paper is to propose a strategy that improves the accuracy of sketching via a boosting procedure that employs a full vector $\tilde{\mathbf{b}}$ corresponding to an inexpensive low-fidelity approximation to $\mathbf{b}$.

Bi-fidelity frameworks assume the availability of a low-fidelity simulation $\tilde{\mathcal{T}}$; that is, a map $\tilde{\mathcal{T}} : \mathbb{R}^q \rightarrow \mathbb{R}$ such that $\tilde{\mathcal{T}}$ is parameterically correlated with $\mathcal{T}$ in some sense, but need not be close to $\mathcal{T}$ in terms of sampled values. Such properties arise, for example, in parametric PDE contexts when $\tilde{\mathcal{T}}$ arises as the discretized PDE solution operator on a spatial mesh that is coarser (and hence less trusted) than the mesh corresponding to $\mathcal{T}$. The decreased accuracy/trustworthiness of $\tilde{\mathcal{T}}$ is balanced by its decreased cost, so that employment of $\tilde{\mathcal{T}}$ may not furnish precise high-fidelity information, but may provide useful knowledge in terms of dependence on the parameter $\mathbf{p}$ with substantially reduced cost.

In the context of constructing our emulator (1.2), our core assumption is that the low-fidelity operator $\tilde{\mathcal{T}}$ is cheap enough so that full exploration of the response over the sampled parameter set $\{\mathbf{p}_i\}_{i \in [N]}$ is more computationally feasible, resulting in a vector $\tilde{\mathbf{b}} \in \mathbb{R}^N$ with low-fidelity entries

$$\tilde{\mathbf{b}}(n) = \sqrt{w_n} \tilde{\mathcal{T}}(\mathbf{p}_n). \quad (2.8)$$

Of course, one may propose constructing the emulator $\mathcal{T}$ in (1.2) by simply replacing $\mathbf{b}$ by $\tilde{\mathbf{b}}$, but this restricts the accuracy of the emulator $\mathcal{T}$ to the potentially bad accuracy of $\tilde{\mathcal{T}}$. In this paper, we propose a more sophisticated use of $\tilde{\mathbf{b}}$, in conjunction with a single sparse sketch of $\mathbf{b}$, that retains some accuracy characteristics of $\mathbf{x}^*$.

3 Bi-fidelity boosting (BFB) in sketched least squares problems

In practice, one often requires the probability of successfully obtaining a good approximation $\mathbf{x}^*$ associated with a random sketch from section 2.2 to be sufficiently close to 1, and one way to achieve this with fixed sketch size is through a boosting procedure. Assuming the availability of a collection of sketching matrices $\{\mathbf{S}_\ell \in \mathbb{R}^{m \times N}\}_{\ell \in [L]}$, one computes the residual for the $\mathbf{S}_\ell$ -sketched solution (i.e., $\|\mathbf{A}(\mathbf{S}_\ell \mathbf{A})^\dagger (\mathbf{S}_\ell \mathbf{b}) - \mathbf{b}\|_2$) for each $\mathbf{S}_\ell$ and then selects the one that yields the smallest residual for use. Even if each sketch is sparse, this straightforward procedure inflates the required sampling cost of the forward model $\mathcal{T}$ by the factor L , which may be computationally prohibitive. To ameliorate this boosting cost, we employ a bi-fidelity strategy.

In Section 3.1 we present our proposed algorithm for quadrature sampling which leverages sketching BFB. Sections 3.2 and 3.3 give our pre-asymptotic and asymptotic analysis results, respectively. We collect some preliminary technical results in section 3.4, and prove our pre-asymptotic results in section 3.5. The asymptotic result is proven in Appendix B. We end with section 3.6 that provides results for random sketches achieving the (ε, δ) condition in Definition 2.1.

3.1 Proposed algorithm

A distinguishing feature of the least squares problem in our setup is that full information of the high-fidelity data $\mathbf{b}$ is unaffordable due to computational restrictions; instead, we can only afford to generate a small number of entries of $\mathbf{b}$. Meanwhile, the low-fidelity data vector $\tilde{\mathbf{b}} \in \mathbb{R}^N$ that exhibits some type of correlation with $\mathbf{b}$ is readily available for repeated use. (This correlation-like condition is quantifying through the parameter ν introduced in Theorem 3.2.) We propose a modified boosting procedure, where the boosting phase of a sketched least squares problem replaces high-fidelity data with low-fidelity data to find the “best” sketching operator and then employs this best sketch directly with high-fidelity data to compute an approximate least squares solution. This procedure is outlined in Algorithm 2.

Algorithm 2: Bi-Fidelity Quadrature Boosting (BFB)

Input: design matrix $\mathbf{A}$, low-fidelity vector $\tilde{\mathbf{b}}$, method for computing entries of the high-fidelity vector $\mathbf{b}$, collection of sketches for boosting $\{\mathbf{S}_\ell\}_{\ell \in [L]}$

Output: an approximate solution $\hat{\mathbf{x}}_{\text{BFB}}$ to (1.2)

- 1: **for** $\ell \in [L]$ **do**
- 2: compute the ℓ -th sketched solution $\hat{\mathbf{x}}_\ell$ using the low-fidelity data:

$$\hat{\mathbf{x}}_\ell = \arg \min_{\mathbf{x} \in \mathbb{R}^d} \|\mathbf{S}_\ell \mathbf{A} \mathbf{x} - \mathbf{S}_\ell \tilde{\mathbf{b}}\|_2$$

- 3: **end for**
- 4: find the best low-fidelity sketch index ℓ^* using boosting:

$$\ell^* = \arg \min_{\ell \in [L]} \|\mathbf{A} \hat{\mathbf{x}}_\ell - \tilde{\mathbf{b}}\|_2$$

- 5: use sketch $\mathbf{S}_{\ell^*}$ to compute an approximate solution to (1.2):

$$\hat{\mathbf{x}}_{\text{BFB}} = \arg \min_{\mathbf{x} \in \mathbb{R}^d} \|\mathbf{S}_{\ell^*} \mathbf{A} \mathbf{x} - \mathbf{S}_{\ell^*} \mathbf{b}\|_2 \quad // \text{ Requires computing } m \text{ entries of } \mathbf{b}$$

The oracle sketch in this scenario is the one identified by the boosting strategy operating directly on the high-fidelity least squares problem, which is computationally unaffordable:

$$\ell^{**} = \arg \min_{\ell \in [L]} \|\mathbf{A} \hat{\mathbf{x}}_\ell - \mathbf{b}\|_2^2, \quad \text{where } \hat{\mathbf{x}}_\ell = \arg \min_{\mathbf{x} \in \mathbb{R}^d} \|\mathbf{S}_\ell \mathbf{A} \mathbf{x} - \mathbf{S}_\ell \tilde{\mathbf{b}}\|_2. \quad (3.1)$$

In the coming sections we will theoretically investigate the sketch transferability between high- and low-fidelity boosting, i.e., when the residual associated to $\hat{\mathbf{x}}_{\text{BFB}}$, the solution produced by Algorithm 2, is comparable to the residual associated to $\hat{\mathbf{x}}_{\ell^{**}}$.

We divide our analysis into two cases: Our first analysis frames performance of Algorithm 2 in terms of an *optimality coefficient*, defined in (3.2), which measures the quality of the least squares residual for a particular sketch $\mathbf{S}$; we provide pre-asymptotic analysis with quantitative results that provides qualitative guidance on how the BFB algorithm behaves in terms of the tradeoff in the number of sketches L versus the optimality coefficient (see the discussion following Theorem 3.4). Our second theoretical result is an asymptotic analysis with Gaussian sketches that confirms the intuition that the probabilistic correlations between the low- and high-fidelity random sketches is high when $\mathbf{b}$ and $\tilde{\mathbf{b}}$ have high vector correlations (see the discussion around Theorem 3.5).

For analysis purposes we make the following assumption.

Assumption 3.1. Assume that neither $\tilde{\mathbf{b}}$ nor $\mathbf{b}$ lie in $\text{range}(\mathbf{A})$, i.e., we assume $\tilde{\mathbf{b}}, \mathbf{b} \notin \text{range}(\mathbf{A})$.

This is a reasonable assumption. If $\mathbf{b} \in \text{range}(\mathbf{A})$, then it would be possible to solve the high-fidelity least squares problem exactly by sampling $m = d$ linearly independent rows of $\mathbf{A}$ and the corresponding rows of $\mathbf{b}$. In this case, it is therefore easy to solve (1.2) and only requires accessing d rows of $\mathbf{b}$. Similarly, if $\tilde{\mathbf{b}} \in \text{range}(\mathbf{A})$ then it would be easy to compute a sketch $\mathbf{S}_\ell$ which only samples $m = d$ rows and achieves zero error in Line 4 of Algorithm 2, therefore making the boosting procedure vacuous.

3.2 Pre-asymptotic analysis via optimality coefficients

We introduce the following measure of relative error difference between the sketched and optimal solutions:

$$\mu_{\mathbf{A}}(\mathbf{b}, \mathbf{S}) \stackrel{\text{def}}{=} \sqrt{\frac{r_{\mathbf{S}}^2(\mathbf{A}, \mathbf{b}) - r^2(\mathbf{A}, \mathbf{b})}{r^2(\mathbf{A}, \mathbf{b})}} \stackrel{(*)}{=} \frac{\|(\mathbf{S}\mathbf{Q})^\dagger \mathbf{S}\mathbf{Q}_\perp \mathbf{Q}_\perp^T \mathbf{b}\|_2}{\|\mathbf{Q}_\perp \mathbf{Q}_\perp^T \mathbf{b}\|_2}, \quad (3.2)$$

where $\mathbf{Q} = \text{orth}(\mathbf{A})$, and the second equality marked $(*)$ is valid if $\text{rank}(\mathbf{S}\mathbf{A}) = \text{rank}(\mathbf{A})$, which we establish in Lemma 3.7. For notational simplicity we usually drop the subscript and write $\mu(\mathbf{b}, \mathbf{S})$ when $\mathbf{A}$ is clear from context, but we emphasize that μ does depend on $\mathbf{A}$. Note that $r(\mathbf{A}, \mathbf{b}) = \|\mathbf{Q}_\perp \mathbf{Q}_\perp^T \mathbf{b}\|_2 > 0$ due to Assumption 3.1, so the denominator in (3.2) is nonzero. We call μ the *optimality coefficient*. Smaller values of μ are better in practice: $\mu = 0$ implies the sketch achieves perfect reconstruction of the data relative to the full least squares solution.

We provide two main theoretical results which shed light on the performance of Algorithm 2 from two different perspectives. The first result shows that with an appropriate choice of the sketches $\{\mathbf{S}_\ell\}_{\ell \in [L]}$, Algorithm 2 produces a solution whose relative error is close to that of the oracle sketch solution in (3.1). Note that it would be straightforward to provide such guarantees if $r_{\mathbf{S}}(\mathbf{A}, \tilde{\mathbf{b}}) \leq r_{\mathbf{S}'}(\mathbf{A}, \tilde{\mathbf{b}})$ implied $r_{\mathbf{S}}(\mathbf{A}, \mathbf{b}) \leq r_{\mathbf{S}'}(\mathbf{A}, \mathbf{b})$, in which case $\ell^* = \ell^{**}$. This may happen, for instance, when $\tilde{\mathbf{b}}$ and $\mathbf{b}$ differ by a scaling. This monotone property of r when replacing $\mathbf{b}$ with $\tilde{\mathbf{b}}$ is unfortunately unlikely to hold in practice. Our result, which appears in Theorem 3.2, identifies alternative conditions that ensure $\mathbf{S}_{\ell^*}$ is a “good” sketch for the high-fidelity data.

Theorem 3.2. *Fix a positive integer L and suppose $\delta, \varepsilon \in (0, 1]$. If $\{\mathbf{S}_\ell\}_{\ell \in [L]}$ is a sequence of i.i.d. random matrices whose distribution is an $(\varepsilon, \frac{\delta}{L})$ pair for $(\mathbf{Q}, \mathbf{h})$, where*

$$\mathbf{h} \stackrel{\text{def}}{=} ((\mathbf{P}_{\mathbf{Q}_\perp} \mathbf{b})_{\mathcal{P}} - (\mathbf{P}_{\mathbf{Q}_\perp} \tilde{\mathbf{b}})_{\mathcal{P}})_{\mathcal{P}} \quad \text{and} \quad \mathbf{Q} \stackrel{\text{def}}{=} \text{orth}(\mathbf{A}),$$

then with probability at least $1 - \delta$,

$$\mu(\mathbf{b}, \mathbf{S}_{\ell^*}) \leq \mu(\mathbf{b}, \mathbf{S}_{\ell^{**}}) + 2\sqrt{6(1 - \nu)}\varepsilon, \quad (3.3)$$

where ν denotes the absolute correlation coefficient between $\mathbf{P}_{\mathbf{Q}_\perp} \mathbf{b}$ and $\mathbf{P}_{\mathbf{Q}_\perp} \tilde{\mathbf{b}}$:

$$\nu \stackrel{\text{def}}{=} |\text{corr}(\mathbf{P}_{\mathbf{Q}_\perp} \mathbf{b}, \mathbf{P}_{\mathbf{Q}_\perp} \tilde{\mathbf{b}})|. \quad (3.4)$$

In addition, on the event where (3.3) is true, we also have that (2.4) holds with $\mathbf{S} = \mathbf{S}_\ell$ for every $\ell \in [L]$.

Theorem 3.2 shows that if a sketch satisfies an $(\varepsilon, \delta/L)$ condition for the pair $\mathbf{Q}$ and an element $\mathbf{h}$ of $\text{range}(\mathbf{Q}_\perp)$, then we are able to prove bounds on the low-fidelity boosted optimality coefficient $\mu(\mathbf{b}, \mathbf{S}_{\ell^*})$ relative to the oracle high-fidelity boosted optimality coefficient $\mu(\mathbf{b}, \mathbf{S}_{\ell^{**}})$. This is quite a general statement that accommodates a wide range of sketching operators. The condition on the operators $\{\mathbf{S}_\ell\}_{\ell \in [L]}$ is, for example, satisfied by all sketching operators in Sections 2.2.2–2.2.4 when the embedding dimension m is sufficiently large. More precise statements for the leverage score and Gaussian sketches are provided in Theorem 3.11.

In order to achieve a good approximate solution when applying sketching techniques in least squares problems the sketching operator must preserve the relevant geometry of the problem. In particular, it is key that $\mathbf{Q}$ and $\mathbf{P}_{\mathbf{Q}_\perp} \mathbf{b}$ remain roughly orthogonal after the sketching operator has been applied. This importance of preserving $\mathbf{P}_{\mathbf{Q}_\perp} \mathbf{b}$ in the sketching phase when $\mathbf{b}$ is replaced by low-fidelity data $\tilde{\mathbf{b}}$ manifests in Theorem 3.2 through the correlation parameter ν .

Remark 3.3. Equation (3.3) suggests that $\mathbf{S}_{\ell^*}$ is “good” when ν is large. This explicitly requires high parametric correlation between the portions of $\mathbf{b}$ and $\tilde{\mathbf{b}}$ that lie orthogonal to the range of $\mathbf{A}$. A more subtle sufficient condition ensuring large ν is furnished by our discussion following Proposition 3.8, which provides a lower bound for ν in terms of other parameters.

Theorem 3.2 does not provide a concrete strategy for how the sketches used in boosting are chosen or constructed. However, near-optimal sketches (in particular satisfying our required (ϵ, δ) pair condition) are known to be produced through the well-known randomized approaches discussed in sections 2.2.2-2.2.4. Precise statements for such sketch estimates are given later in by Theorem 3.11 in section 3.6, but it is appropriate for us to establish here that combining Theorem 3.2 with good sketching techniques results in explicit and illuminating theory for Algorithm 2. In particular, one expects a tradeoff between the values of ν and L : boosting with a large number L of sketches should work up to a threshold determined by the amount of correlation between $\mathbf{b}$ and $\tilde{\mathbf{b}}$. I.e., any accuracy gained by BFB should be limited by how correlated the low- and high-fidelity models are, and one expects this to manifest in a relationship between L and ν . The theory we develop below reveals this tradeoff. We focus on generating the sketches $\{\mathbf{S}_{\ell}\}_{\ell \in [L]}$ through leverage score sampling, as described explicitly by (2.6) in section 2.2.2. We briefly discuss afterward that one could generalize the result to more general sketches.

Theorem 3.4. *Let $\delta, \epsilon \in (0, 1/2)$ and $L \in \mathbb{N}$ be chosen, and assume*

$$d \leq \frac{\delta}{4} \exp\left(\frac{2}{35\epsilon\delta}\right). \quad (3.5)$$

Now consider Algorithm 2, where $\{\mathbf{S}_{\ell}\}_{\ell \in [L]}$ are iid samples of a leverage score sketching operator defined in (2.6), with the sampling requirement

$$m \geq \frac{4dL}{\epsilon\delta}. \quad (3.6)$$

Then each $\mathbf{S}_{\ell}$ satisfies an $(\epsilon/L, \delta/2)$ condition for the pair $(\mathbf{Q}, \mathbf{h})$, and with probability at least $1 - \delta$, we have

$$r_{\mathbf{S}_{\ell^*}}^2(\mathbf{A}, \mathbf{b}) \leq \left[1 + \frac{\epsilon}{L}\tau\right] r^2(\mathbf{A}, \mathbf{b}), \quad (3.7)$$

where

$$\tau = \tau(\epsilon, \delta, \nu, L) = 24L(1 - \nu) + \frac{\delta}{2} \left(1 + 4\sqrt{6(1 - \nu)\epsilon}\right).$$

The results above give explicit behavior of the BFB residual via a concrete sketching strategy for Algorithm 2. Note in particular that the sampling requirement $m = \mathcal{O}(L/\epsilon)$ in (3.6) means that *without* boosting and simply generating one sketch $\mathbf{S}$ according to (3.6), which requires m high-fidelity samples (equivalent to the number from BFB), we expect that the residual from this one sketch behaves like

$$r_{\mathbf{S}}^2(\mathbf{A}, \mathbf{b}) \sim \left(1 + \frac{\epsilon}{L}\right) r^2(\mathbf{A}, \mathbf{b}).$$

Comparing the above to (3.7), note that the only difference is the appearance of τ , and hence we expect BFB to be useful (compared to an equivalent number of high-fidelity samples devoted to a non-boosting strategy) when $\tau \leq 1$, which requires,

$$L \lesssim \frac{1}{1 - \nu}.$$

I.e., boosting with L sketches is useful in BFB up to a threshold $\sim 1/(1 - \nu)$. Boosting with *more* than this threshold level of sketches causes the error bound to saturate at a level determined by $1 - \nu$. Since ν is the correlation between the range($\mathbf{A}$)-orthogonal components of $\mathbf{b}$ and $\tilde{\mathbf{b}}$, we conclude that highly correlated

range-orthogonal residuals (large values of ν very close to 1) are optimal for BFB in the sense that sketching with large L will be effective.

A second observation we make is that the $m \sim L$ requirement (3.6) is theoretically suboptimal. In particular, we show in Theorem 3.11 that stronger coherence-like conditions on the matrix $\mathbf{A}$ imply that leverage score sketching with $m \sim \log L$ is sufficient to achieve the requisite $(\epsilon/L, \delta)$ condition, see (3.26) in Theorem 3.11. We also note that Gaussian sketches only require $m \sim \log L$ samples (see (3.23)), and one can achieve the (ϵ, δ) condition *on average* using $m \sim \log L$ samples (see, e.g., [Mal+22, Equation (2.18)]). Finally, if (3.5) is violated, then indeed $m \sim \log L$ (see (3.24) and the intermediate computation in (3.8)) for leverage score sketches. Thus, we expect in practice that $m \sim \log L$ samples are sufficient.

We give the proof of theorem 3.4 below to demonstrate how it relies on Theorem 3.2; we will prove Theorem 3.2 in the coming sections.

Proof of Theorem 3.4. We start by making two conclusions from the conditions (3.5) and (3.6). First, under these conditions,

$$35 \log \left(\frac{4d}{\delta} \right) \leq \frac{2}{\epsilon \delta} \implies m \geq d \max \left\{ 35 \log \left(\frac{4dL}{(\delta/2)} \right), \frac{2L}{\epsilon(\delta/2)} \right\}, \quad (3.8)$$

implying that condition (3.24) holds, so that result 2 from Theorem 3.11 guarantees that the distribution from which the $\mathbf{S}_\ell$ sketches are drawn satisfies and $(\epsilon, \frac{\delta}{2L})$ condition. Thus, theorem 3.2 states that there is an event E_1 such that

$$\Pr(E_1) \geq 1 - \delta/2, \quad \text{On event } E_1, \text{ then (3.3) holds.} \quad (3.9)$$

The above is our first conclusion. For our second conclusion, we note that (3.6) and (3.5) imply,

$$m \geq \frac{2d}{\left[\frac{\epsilon}{L} \left(\frac{\delta}{2} \right)^{1-1/L} \right] \left(\frac{\delta}{2} \right)^{1/L}},$$

so that again we satisfy (3.24) (employing a variation of the argument (3.8)), and so by Theorem 3.11, the distribution from which $\mathbf{S}_\ell$ is drawn satisfies an $(\tilde{\epsilon}, \tilde{\delta})$ condition for $(\mathbf{A}, \mathbf{b})$, where,

$$\tilde{\epsilon} \stackrel{\text{def}}{=} \frac{\epsilon}{L} \left(\frac{\delta}{2} \right)^{1-1/L}, \quad \tilde{\delta} \stackrel{\text{def}}{=} \left(\frac{\delta}{2} \right)^{1/L}.$$

Therefore with probability at least $1 - \tilde{\delta}$,

$$r_{\mathbf{S}_\ell}^2(\mathbf{A}, \mathbf{b}) \leq (1 + \tilde{\epsilon})r^2(\mathbf{A}, \mathbf{b}),$$

so that a union bound implies that there is an event E_2 on which our second conclusion holds:

$$\Pr(E_2) \geq 1 - (\tilde{\delta})^L = 1 - \delta/2 \quad \text{On event } E_2, \text{ then } \min_{\ell \in [L]} r_{\mathbf{S}_\ell}^2(\mathbf{A}, \mathbf{b}) \leq (1 + \tilde{\epsilon})r_{\mathbf{S}_{\ell^{**}}}^2(\mathbf{A}, \mathbf{b}). \quad (3.10)$$

We now observe that for any $\eta > 0$, the bound

$$|\mu(\mathbf{b}, \mathbf{S}_{\ell^*}) - \mu(\mathbf{b}, \mathbf{S}_{\ell^{**}})| \leq \eta$$

implies that

$$r_{\mathbf{S}_{\ell^*}}^2(\mathbf{A}, \mathbf{b}) \leq r_{\mathbf{S}_{\ell^{**}}}^2(\mathbf{A}, \mathbf{b}) + r^2(\mathbf{A}, \mathbf{b})(\eta^2 + 2\eta\mu(\mathbf{b}, \mathbf{S}_{\ell^{**}})).$$

Thus, $E_1 \cap E_2$ occurs with probability at least $1 - \delta$, and on this event (3.9) ensures that η is given by the right-hand side of (3.3). Also, on this event (3.10) implies that $\mu(\mathbf{b}, \mathbf{S}_{\ell^{**}}) = \tilde{\epsilon}$, i.e., $r_{\mathbf{S}_{\ell^{**}}}^2(\mathbf{A}, \mathbf{b}) \leq (1 + \tilde{\epsilon})r^2(\mathbf{A}, \mathbf{b})$. Using these expressions in the above inequality, simplifying, and using $(\delta/2)^{1-1/L} \leq \delta/2$ yields the result (3.7). $\square$

We emphasize that the proof above shows how Theorem 3.2 can be used to prove results like Theorem 3.4 for more general sketches.

3.3 Asymptotic analysis via probabilistic correlation

We provide alternative analysis of Algorithm 2 motivated by the following intuition: If $\mu(\mathbf{b}, \mathbf{S})$ and $\mu(\tilde{\mathbf{b}}, \mathbf{S})$ are probabilistically correlated in some sense, then we expect that Algorithm 2 should produce a sketching operator $\mathbf{S}_{\ell^*}$ that is close to the oracle sketch $\mathbf{S}_{\ell^{**}}$. We give a technical verification of this intuition below in Theorem 3.5, providing an asymptotic lower bound on a certain measure of correlation between the two optimality coefficients when $\mathbf{S}$ is a Gaussian sketching operator.

Theorem 3.5. *If $\mathbf{S}$ is a Gaussian sketch, then*

$$\liminf_{m \rightarrow \infty} \text{corr}(\mu^2(\mathbf{b}, \mathbf{S}), \mu^2(\tilde{\mathbf{b}}, \mathbf{S})) \geq \frac{\|\mathbf{P}_{\mathbf{Q}_\perp} \mathbf{b}_{\mathcal{P}}\|_2^2 - \sqrt{6} \min\{\|\mathbf{P}_{\mathbf{Q}_\perp}(\mathbf{b}_{\mathcal{P}} \pm \tilde{\mathbf{b}}_{\mathcal{P}})\|_2\}}{\|\mathbf{P}_{\mathbf{Q}_\perp} \tilde{\mathbf{b}}_{\mathcal{P}}\|_2^2}, \quad (3.11)$$

where $\mathbf{b}_{\mathcal{P}}, \tilde{\mathbf{b}}_{\mathcal{P}}$ are normalized versions of $\mathbf{b}$ and $\tilde{\mathbf{b}}$, respectively, and the minimum is taken over the two $\pm$ options. Moreover, if

$$\varphi \stackrel{\text{def}}{=} \frac{|\langle \mathbf{b}, \tilde{\mathbf{b}} \rangle|}{\|\mathbf{b}\|_2 \|\tilde{\mathbf{b}}\|_2} \geq \frac{\|\mathbf{P}_{\mathbf{Q}} \mathbf{b}\|_2}{\|\mathbf{b}\|_2} \stackrel{\text{def}}{=} \kappa, \quad (3.12)$$

then we further have that

$$\liminf_{m \rightarrow \infty} \text{corr}(\mu^2(\mathbf{b}, \mathbf{S}), \mu^2(\tilde{\mathbf{b}}, \mathbf{S})) \geq (1 - \kappa^2) - \frac{\sqrt{12(1 - \varphi)}}{(\varphi - \kappa)^2}. \quad (3.13)$$

In Theorem 3.5 we restrict to Gaussian sketches and consider $\text{corr}(\mu^2(\mathbf{b}, \mathbf{S}), \mu^2(\tilde{\mathbf{b}}, \mathbf{S}))$ (rather than the more natural quantity $\text{corr}(\mu(\mathbf{b}, \mathbf{S}), \mu(\tilde{\mathbf{b}}, \mathbf{S}))$) in order to make analysis tractable. In general $\text{corr}(\mu(\mathbf{b}, \mathbf{S}), \mu(\tilde{\mathbf{b}}, \mathbf{S}))$ and $\text{corr}(\mu^2(\mathbf{b}, \mathbf{S}), \mu^2(\tilde{\mathbf{b}}, \mathbf{S}))$ may have significantly different statistical properties. However, if either of them is close to 1, then that would indicate a monotonically increasing (although not necessarily linear) relationship between $\mu(\mathbf{b}, \mathbf{S})$ and $\mu(\tilde{\mathbf{b}}, \mathbf{S})$, and when such a relationship holds we expect the boosting procedure in Algorithm 2 to work well. While we restrict to Gaussian sketches, this probabilistic model is usually a good indicator of how other sketches perform [MT20, Remark 8.2]. I.e., we expect the result to carry over to the random sampling-based sketches (e.g., leverage scores) that we consider. We verify this numerically in Section 4.

Remark 3.6. The lower bound in (3.13) is useful only when the right-hand side is close to 1, which roughly requires φ to be large and κ to be small. See Remark 3.9 for how this condition relates to Theorem 3.2.

The rest of this section is organized as follows. Section 3.4 derives some preliminary technical results. Section 3.5 then proves Theorem 3.2. Section 3.6 provides theoretical guarantees for when various sketches satisfy the (ε, δ) pair condition in Definition 2.1 and discuss how this condition in turn ensures that those sketching operators satisfy the requirements in Theorem 3.2. The proof of Theorem 3.5 is given in Appendix B.

3.4 Preliminary technical results

Our first task is to understand how the optimal residual $r(\mathbf{A}, \mathbf{b})$ compares to $r_{\mathbf{S}}(\mathbf{A}, \mathbf{b})$. Throughout this section let $\mathbf{Q} = \text{orth}(\mathbf{A})$.

Lemma 3.7. *Given a sketch matrix $\mathbf{S}$, assume $\ker(\mathbf{S}) \cap \text{range}(\mathbf{A}) = \{\mathbf{0}\}$, or, equivalently, $\text{rank}(\mathbf{S}\mathbf{A}) = \text{rank}(\mathbf{A})$. Then we have,*

$$r_{\mathbf{S}}^2(\mathbf{A}, \mathbf{b}) = r^2(\mathbf{A}, \mathbf{b}) + \|(\mathbf{S}\mathbf{Q})^\dagger \mathbf{S}\mathbf{Q}_\perp \mathbf{Q}_\perp^T \mathbf{b}\|_2^2.$$

Proof. Under the assumption $\ker(\mathbf{S}) \cap \text{range}(\mathbf{A}) = \{\mathbf{0}\}$, the sketched least squares problem reproduces elements of $\text{range}(\mathbf{A})$: For any $\mathbf{c} \in \text{range}(\mathbf{A})$,

$$\mathbf{A}(\mathbf{S}\mathbf{A})^\dagger \mathbf{S}\mathbf{c} = \mathbf{c}. \quad (3.14)$$

The solution to the sketched least squares problem (2.2) is $(\mathbf{S}\mathbf{A})^\dagger \mathbf{S}\mathbf{b}$. Combining this fact with (2.3) and (3.14) yields

$$r_{\mathbf{S}}^2(\mathbf{A}, \mathbf{b}) = \|\mathbf{b} - \mathbf{A}(\mathbf{S}\mathbf{A})^\dagger \mathbf{S}\mathbf{b}\|_2^2 = \|\mathbf{b} - \mathbf{A}(\mathbf{S}\mathbf{A})^\dagger \mathbf{S}(\mathbf{Q}\mathbf{Q}^T + \mathbf{Q}_\perp \mathbf{Q}_\perp^T) \mathbf{b}\|_2^2 = r^2(\mathbf{A}, \mathbf{b}) + \|(\mathbf{S}\mathbf{Q})^\dagger \mathbf{S}\mathbf{Q}_\perp \mathbf{Q}_\perp^T \mathbf{b}\|_2^2.$$

□

We conclude that $r_{\mathbf{S}}(\mathbf{A}, \mathbf{b})$ is comparable to $r(\mathbf{A}, \mathbf{b})$ if and only if $\|(\mathbf{S}\mathbf{Q})^\dagger \mathbf{S}\mathbf{Q}_\perp \mathbf{Q}_\perp^T \mathbf{b}\|_2^2$ is small. The quantities ν , φ and κ defined in (3.4) and (3.12) are related by the following inequality.

Proposition 3.8. *Assume $\varphi \geq \kappa$. Then we have the two inequalities,*

$$\nu \geq \varphi - \kappa \min \left\{ 1, \sqrt{2(1 - \varphi + \kappa)} \right\}. \quad (3.15)$$

$$\nu \geq \varphi - (\varphi \tilde{\kappa} + \sqrt{1 - \varphi^2}) \min \left\{ 1, \sqrt{2(1 - \varphi + \varphi \tilde{\kappa} + \sqrt{1 - \varphi^2})} \right\}. \quad (3.16)$$

where

$$\tilde{\kappa} \stackrel{\text{def}}{=} \frac{\|\mathbf{P}_Q \tilde{\mathbf{b}}\|_2}{\|\tilde{\mathbf{b}}\|_2},$$

measures the relative energy of the low-fidelity vector in the range of $\mathbf{A}$.

Proof. We first prove (3.15). Since correlation coefficients are scale-invariant, without loss of generality we assume $\|\mathbf{b}\|_2 = \|\tilde{\mathbf{b}}\|_2 = 1$. Write down the orthogonal decomposition of $\mathbf{b}$ and $\tilde{\mathbf{b}}$ in $\mathbf{Q} \oplus \mathbf{Q}_\perp$ as follows:

$$\begin{aligned} \mathbf{b} &= \underbrace{\mathbf{P}_Q \mathbf{b}}_{\mathbf{b}_1} + \underbrace{\mathbf{P}_{Q_\perp} \mathbf{b}}_{\mathbf{b}_2}, \\ \tilde{\mathbf{b}} &= \underbrace{\mathbf{P}_Q \tilde{\mathbf{b}}}_{\tilde{\mathbf{b}}_1} + \underbrace{\mathbf{P}_{Q_\perp} \tilde{\mathbf{b}}}_{\tilde{\mathbf{b}}_2}. \end{aligned}$$

Notice that $\|\mathbf{b}_1\|_2^2 + \|\mathbf{b}_2\|_2^2 = \|\tilde{\mathbf{b}}_1\|_2^2 + \|\tilde{\mathbf{b}}_2\|_2^2 = 1$. It follows from the Cauchy–Schwarz inequality and the definitions in (3.4) and (3.12) that

$$\nu = \frac{|\langle \mathbf{b}_2, \tilde{\mathbf{b}}_2 \rangle|}{\|\mathbf{b}_2\|_2 \|\tilde{\mathbf{b}}_2\|_2} \geq |\langle \mathbf{b}, \tilde{\mathbf{b}} \rangle - \langle \mathbf{b}_1, \tilde{\mathbf{b}}_1 \rangle| \geq \varphi - \|\mathbf{b}_1\|_2 \|\tilde{\mathbf{b}}_1\|_2 = \varphi - \kappa \|\tilde{\mathbf{b}}_1\|_2 \geq \varphi - \kappa. \quad (3.17)$$

The last inequality can be replaced by a more accurate estimate for $\|\tilde{\mathbf{b}}_1\|_2$:

$$\varphi = |\langle \mathbf{b}, \tilde{\mathbf{b}} \rangle| = |\langle \mathbf{b}_1, \tilde{\mathbf{b}}_1 \rangle + \langle \mathbf{b}_2, \tilde{\mathbf{b}}_2 \rangle| \leq \|\mathbf{b}_1\|_2 \|\tilde{\mathbf{b}}_1\|_2 + \|\mathbf{b}_2\|_2 \|\tilde{\mathbf{b}}_2\|_2 \leq \kappa + \|\tilde{\mathbf{b}}_2\|_2 = \sqrt{1 - \|\tilde{\mathbf{b}}_1\|_2^2} + \kappa, \quad (3.18)$$

which can be reorganized as

$$\|\tilde{\mathbf{b}}_1\|_2 \leq \sqrt{1 - (\varphi - \kappa)^2} = \sqrt{(1 - \varphi + \kappa)(1 + \varphi - \kappa)} \leq \sqrt{2(1 - \varphi + \kappa)}. \quad (3.19)$$

Combining (3.17) and (3.19) finishes the proof of (3.15).

To show (3.16), we again assume $\|\mathbf{b}\|_2 = \|\tilde{\mathbf{b}}\|_2 = 1$, so that,

$$\kappa = \|\mathbf{P}_Q \mathbf{b}\|_2 = \|\mathbf{P}_Q (\mathbf{P}_{\tilde{\mathbf{b}}} \mathbf{b} + \mathbf{b} - \mathbf{P}_{\tilde{\mathbf{b}}} \mathbf{b})\|_2 \leq \varphi \|\mathbf{P}_Q \tilde{\mathbf{b}}\|_2 + \|\mathbf{b} - \mathbf{P}_{\tilde{\mathbf{b}}} \mathbf{b}\|_2 = \varphi \tilde{\kappa} + \sqrt{1 - \varphi^2}.$$

Plugging this into (3.15) and noting that the right-hand side of (3.15) is decreasing in κ yields (3.16). □

The main appeal of (3.16) is that the quantity $\tilde{\kappa}$ involves only low-fidelity data, and hence can be estimated. I.e., (3.16) gives a more practically computable lower bound for ν , involving one quantity $\tilde{\kappa}$ that depends only on low-fidelity data $\tilde{\mathbf{b}}$, and the correlation φ between $\mathbf{b}$ and $\tilde{\mathbf{b}}$.

Remark 3.9. Recall that our main convergence result, Theorem 3.2, has more attractive bounds when ν is large. By (3.15), ν is large if $\varphi \approx 1$ and $\varphi \gg \kappa$, which coincides with sufficient conditions to ensure attractive bounds in (3.13) in Theorem 3.5. (Cf. Remark 3.6.) Thus, $\varphi \gg \kappa$ is a unifying condition under which both of our main theoretical results, Theorem 3.2 and Theorem 3.5, provide useful bounds. The condition $\varphi \gg \kappa$ means that the correlation between $\mathbf{b}$ and $\tilde{\mathbf{b}}$ is high and strongly dominates the relative energy of $\mathbf{b}$ in $\text{range}(\mathbf{A})$. This condition may seem counterintuitive as it requires the high-fidelity solution to have a relatively large residual. Since μ is defined relative to $r(\mathbf{A}, \mathbf{b})$, a small $r_{\mathbf{S}_{\ell^*}}(\mathbf{A}, \mathbf{b})$ may still result in a large $\mu(\mathbf{b}, \mathbf{S}_{\ell^*})$ even if $r_{\mathbf{S}_{\ell^*}}(\mathbf{A}, \mathbf{b})$ is small but relatively large compared to $r(\mathbf{A}, \mathbf{b})$.

3.5 Proof of Theorem 3.2

We first consider the case $\text{corr}(\mathbf{P}_{\mathbf{Q}_{\perp}} \mathbf{b}, \mathbf{P}_{\mathbf{Q}_{\perp}} \tilde{\mathbf{b}}) \geq 0$. Fixing $\ell \in [L]$, $\mathbf{S} = \mathbf{S}_{\ell}$, consider the event E of probability at least $1 - \delta/L$ where the rank condition in (2.4) holds. On this event, this rank condition with Lemma 3.7 implies that,

$$r_{\mathbf{S}}^2(\mathbf{A}, \mathbf{b}) - r^2(\mathbf{A}, \mathbf{b}) = \|(\mathbf{S}\mathbf{Q})^{\dagger} \mathbf{S}\mathbf{Q}_{\perp} \mathbf{Q}_{\perp}^T \mathbf{b}\|_2^2,$$

allowing us to directly estimate the difference between $\mu(\mathbf{b}, \mathbf{S})$ and $\mu(\tilde{\mathbf{b}}, \mathbf{S})$ as follows:

$$\begin{aligned} |\mu(\mathbf{b}, \mathbf{S}) - \mu(\tilde{\mathbf{b}}, \mathbf{S})| &= \left| \frac{\|(\mathbf{S}\mathbf{Q})^{\dagger} \mathbf{S}\mathbf{Q}_{\perp} \mathbf{Q}_{\perp}^T \mathbf{b}\|_2}{\|\mathbf{Q}_{\perp} \mathbf{Q}_{\perp}^T \mathbf{b}\|_2} - \frac{\|(\mathbf{S}\mathbf{Q})^{\dagger} \mathbf{S}\mathbf{Q}_{\perp} \mathbf{Q}_{\perp}^T \tilde{\mathbf{b}}\|_2}{\|\mathbf{Q}_{\perp} \mathbf{Q}_{\perp}^T \tilde{\mathbf{b}}\|_2} \right| \\ &\leq \|(\mathbf{S}\mathbf{Q})^{\dagger} \mathbf{S} ((\mathbf{P}_{\mathbf{Q}_{\perp}} \mathbf{b})_{\mathcal{P}} - (\mathbf{P}_{\mathbf{Q}_{\perp}} \tilde{\mathbf{b}})_{\mathcal{P}})\|_2 \\ &= \|(\mathbf{P}_{\mathbf{Q}_{\perp}} \mathbf{b})_{\mathcal{P}} - (\mathbf{P}_{\mathbf{Q}_{\perp}} \tilde{\mathbf{b}})_{\mathcal{P}}\|_2 \|(\mathbf{S}\mathbf{Q})^{\dagger} \mathbf{S} \mathbf{h}\|_2 \\ &= \sqrt{\|(\mathbf{P}_{\mathbf{Q}_{\perp}} \mathbf{b})_{\mathcal{P}}\|_2^2 + \|(\mathbf{P}_{\mathbf{Q}_{\perp}} \tilde{\mathbf{b}})_{\mathcal{P}}\|_2^2 - 2\langle (\mathbf{P}_{\mathbf{Q}_{\perp}} \mathbf{b})_{\mathcal{P}}, (\mathbf{P}_{\mathbf{Q}_{\perp}} \tilde{\mathbf{b}})_{\mathcal{P}} \rangle} \|(\mathbf{S}\mathbf{Q})^{\dagger} \mathbf{S} \mathbf{h}\|_2 \\ &= \sqrt{2 - 2\nu} \cdot \|(\mathbf{S}\mathbf{Q})^{\dagger} \mathbf{S} \mathbf{h}\|_2 \\ &= \sqrt{2 - 2\nu} \cdot \|\mathbf{Q}(\mathbf{S}\mathbf{Q})^{\dagger} \mathbf{S} \mathbf{h}\|_2, \end{aligned} \tag{3.20}$$

where the first inequality follows from the reverse triangle inequality, the second to last equality follows (3.4), and the final equality follows from unitary invariance of the operator norm. The case $\text{corr}(\mathbf{P}_{\mathbf{Q}_{\perp}} \mathbf{b}, \mathbf{P}_{\mathbf{Q}_{\perp}} \tilde{\mathbf{b}}) < 0$ can be treated similarly by noting that the inequality on the second line of (3.20) still holds if the minus sign on the right-hand side is changed to a plus sign. The rest of the computation is then done similarly to the case with non-negative correlation.

Note that $(\mathbf{S}\mathbf{Q})^{\dagger} \mathbf{S} \mathbf{h}$ is the $\mathbf{S}$ -sketched least squares solution to $\min_{\mathbf{x}} \|\mathbf{Q}\mathbf{x} - \mathbf{h}\|_2$. Also, note that $\mathbf{h} \in \text{range}(\mathbf{Q}_{\perp})$. Using the residual bound in (2.4), the following also holds on our probabilistic event E :

$$\|\mathbf{Q}(\mathbf{S}\mathbf{Q})^{\dagger} \mathbf{S} \mathbf{h}\|_2^2 + \|\mathbf{h}\|_2^2 = \|\mathbf{Q}(\mathbf{S}\mathbf{Q})^{\dagger} \mathbf{S} \mathbf{h} - \mathbf{h}\|_2^2 \leq (1 + \varepsilon)^2 \min_{\mathbf{x} \in \mathbb{R}^d} \|\mathbf{Q}\mathbf{x} - \mathbf{h}\|_2^2 = (1 + \varepsilon)^2 \|\mathbf{h}\|_2^2.$$

Rearranging terms and noting $\|\mathbf{h}\|_2 = 1$ yields $\|\mathbf{Q}(\mathbf{S}\mathbf{Q})^{\dagger} \mathbf{S} \mathbf{h}\|_2 \leq \sqrt{3\varepsilon}$, which is substituted into (3.20), implying that on an event E with probability at least $1 - \delta/L$, we have

$$|\mu(\mathbf{b}, \mathbf{S}) - \mu(\tilde{\mathbf{b}}, \mathbf{S})| \leq \sqrt{6(1 - \nu)\varepsilon}.$$

Taking a union bound over $\ell \in [L]$ yields that, with probability at least $1 - \delta$,

$$\max_{\ell \in [L]} |\mu(\mathbf{b}, \mathbf{S}_{\ell}) - \mu(\tilde{\mathbf{b}}, \mathbf{S}_{\ell})| \leq \sqrt{6(1 - \nu)\varepsilon}. \tag{3.21}$$

Conditioning on the probabilistic event in (3.21) and using the definition of ℓ^* and ℓ^{**} finishes the proof:

$$\mu(\mathbf{b}, \mathbf{S}_{\ell^*}) \leq \mu(\tilde{\mathbf{b}}, \mathbf{S}_{\ell^*}) + \sqrt{6(1 - \nu)\varepsilon} \leq \mu(\tilde{\mathbf{b}}, \mathbf{S}_{\ell^{**}}) + \sqrt{6(1 - \nu)\varepsilon} \leq \mu(\mathbf{b}, \mathbf{S}_{\ell^{**}}) + 2\sqrt{6(1 - \nu)\varepsilon}.$$

3.6 Achieving the (ε, δ) pair condition

We next show that, for a variety of random sketches of interest, the $(\varepsilon, \frac{\delta}{L})$ pair condition for $(\mathbf{Q}, \mathbf{h})$ in Theorem 3.2 holds for sufficiently large m . We begin with a lemma that gives a sufficient condition for verification of the $(\varepsilon, \frac{\delta}{L})$ pair condition for $(\mathbf{Q}, \mathbf{h})$, which can be deduced as a special case from [Dri+11, Lemma 1]:

Lemma 3.10 (Drineas et al. [Dri+11]). *Let $\mathbf{Q}$ and $\mathbf{h}$ be defined as in Theorem 3.2. The distribution of $\mathbf{S}$ is an $(\varepsilon, \frac{\delta}{L})$ pair for $(\mathbf{Q}, \mathbf{h})$ if the following two conditions hold simultaneously with probability at least $1 - \delta/L$:*

$$\sigma_{\min}^2(\mathbf{S}\mathbf{Q}) \geq \frac{\sqrt{2}}{2} \quad \text{and} \quad \|\mathbf{Q}^T \mathbf{S}^T \mathbf{S} \mathbf{h}\|_2^2 \leq \frac{\varepsilon}{2}, \quad (3.22)$$

where $\sigma_{\min}(\cdot)$ denotes the smallest singular value of a matrix.

When the conditions in Lemma 3.10 hold, one can directly bound (3.20) using the submultiplicativity of operator norms instead of resorting to an (ε, δ) argument as in the proof of Theorem 3.2, although the latter is more general. Theorem 3.11 presents constructive strategies for generating sketch distributions – based on sub-Gaussian random variables and leverage scores – that achieve appropriate (ε, δ) pair conditions. We recall that a random variable X is called sub-Gaussian if, for some $K > 0$ we have $\mathbb{E} \exp(X^2/K^2) \leq 2$ [Ver18, Def. 2.5.6]. The sub-Gaussian norm of X is defined as $\|X\|_{\psi_2} \stackrel{\text{def}}{=} \inf \{K > 0 : \mathbb{E} \exp(X^2/K^2) \leq 2\}$ [Ver18]. A proof of Theorem 3.11 is given in Appendix C. Variants of these results have appeared previously in the literature [DMM06; DMM08; Dri+11; LK20].

Theorem 3.11. *Let $\mathbf{Q}$ and $\mathbf{h}$ be defined as in Theorem 3.2. Write $\mathbf{Q}$ and $\mathbf{S}^T$ as column vectors:*

$$\mathbf{Q} = [\mathbf{q}_1, \dots, \mathbf{q}_d], \quad \mathbf{S}^T = [\mathbf{s}_1, \dots, \mathbf{s}_m],$$

and denote by $q_{ij} \stackrel{\text{def}}{=} \mathbf{q}_i(j)$ and $h_j \stackrel{\text{def}}{=} \mathbf{h}(j)$ the j -th component of $\mathbf{q}_i$ and $\mathbf{h}$, respectively.

1. Suppose $\mathbf{S} \in \mathbb{R}^{m \times N}$ is a dense sketch whose entries are i.i.d. sub-Gaussian random variables with mean 0 and variance $1/m$. Assume the sub-Gaussian norm of each entry of $\sqrt{m}\mathbf{S}$ is bounded by $K \geq 1$. Then the distribution of $\mathbf{S}$ is an $(\varepsilon, \frac{\delta}{L})$ pair for $(\mathbf{Q}, \mathbf{h})$ if

$$m \geq \frac{CK^4}{\varepsilon} d \log \left(\frac{4dL}{\delta} \right), \quad (3.23)$$

where C is an absolute constant.

2. Suppose $\mathbf{S} \in \mathbb{R}^{m \times N}$ is a row sketch based on the leverage scores of $\mathbf{A}$, and $0 < \varepsilon, \delta < 1/2$; see Equation (2.6). Then the distribution of $\mathbf{S}$ is an $(\varepsilon, \frac{\delta}{L})$ pair for $(\mathbf{Q}, \mathbf{h})$ if

$$m \geq \max \left\{ 35d \log \left(\frac{4dL}{\delta} \right), \frac{2dL}{\varepsilon\delta} \right\}. \quad (3.24)$$

Moreover, if

$$\max_{i \in [d]} \max_{j \in [N]: \ell_j > 0} \frac{d|q_{ij}h_j|}{\ell_j} \leq C, \quad \ell_j = \sum_{k \in [d]} q_{kj}^2 \quad (3.25)$$

for some constant $C > 0$, then the distribution of $\mathbf{S}$ is an $(\varepsilon, \frac{\delta}{L})$ pair for $(\mathbf{Q}, \mathbf{h})$ if

$$m \geq \max \left\{ 35, \frac{4C^2}{\varepsilon} \right\} d \log \left(\frac{4dL}{\delta} \right). \quad (3.26)$$

The scalar ℓ_j in (3.25) is the leverage score associated to row j of $\mathbf{A}$, and $(\ell_j)_{j \in [N]}$ defines a (discrete) probability distribution over the row indices $[N]$ of $\mathbf{A}$; see (2.6).

Remark 3.12. When $\mathbf{Q}$ is incoherent, i.e., when its leverage scores satisfy $\ell_i = \mathcal{O}(d/N)$, the entries q_{ij} satisfy $q_{ij} = \mathcal{O}(1/\sqrt{N})$. For any $\mathbf{h}$ such that $\max_{j \in [N]} |h_j| \lesssim \mathcal{O}(1/\sqrt{N})$, the condition in (3.25) is satisfied with $C = \mathcal{O}(1)$:

$$\max_{i \in [d]} \max_{j \in [N]: \ell_j > 0} \frac{d|q_{ij}h_j|}{\ell_j} \lesssim \frac{d \cdot \frac{1}{\sqrt{N}} \cdot \frac{1}{\sqrt{N}}}{\frac{d}{N}} = 1.$$

Remark 3.13. As noted in Section 2.2.3, leveraged volume sampling requires $m \gtrsim d \log(d/\delta) + d/(\varepsilon\delta)$ samples to satisfy the (ε, δ) pair condition. This result appears in Corollary 10 of [DWH18].

4 Numerical experiments

In this section we illustrate various aspects of the BFB approach using both manufactured data as well as data obtained from PDE solutions. The codes used to generate the results of this section are available from the GitHub repository <https://github.com/CU-UQ/BF-Boosted-Quadrature-Sampling>.

4.1 Verification of theoretical results on synthetic data

We first verify the theoretical results in Theorems 3.2 and 3.5. We do this by simulating different values for $\mathbf{S}$, $\mathbf{b}$ and $\tilde{\mathbf{b}}$. We generate a design matrix $\mathbf{A} \in \mathbb{R}^{1000 \times 50}$ (i.e., $N = 1000$ and $d = 50$) with i.i.d. standard normal entries and fix it in the rest of the simulations. For sketching matrices $\mathbf{S}$, we choose the embedding dimension to be $m = 100$ and consider both the Gaussian and leverage score sampling sketches. We generate multiple different versions of the vectors $\mathbf{b}$ and $\tilde{\mathbf{b}}$ that correspond to different values of κ and φ . Recall that these parameters control how much of $\mathbf{b}$ is in the range of $\mathbf{A}$ and the absolute value of the correlation between $\mathbf{b}$ and $\tilde{\mathbf{b}}$, respectively. The vectors are generated via

$$\begin{aligned}\mathbf{b} &= \kappa \mathbf{Q} \mathbf{z}_1 + \sqrt{1 - \kappa^2} \mathbf{Q}_\perp \mathbf{z}_2, \\ \tilde{\mathbf{b}} &= \varphi \mathbf{b} + \sqrt{1 - \varphi^2} \mathbf{b}_\perp \mathbf{z}_3,\end{aligned}$$

where $\mathbf{Q} = \text{orth}(\mathbf{A})$, and $\mathbf{z}_1 \in \mathbb{R}^{d-1}$, $\mathbf{z}_2 \in \mathbb{R}^{N-d-1}$ and $\mathbf{z}_3 \in \mathbb{R}^{N-2}$ are generated by normalizing random vectors of appropriate length whose entries are i.i.d. standard normal. In the experiment, the vectors $\mathbf{z}_1, \mathbf{z}_2, \mathbf{z}_3$ are drawn once and then kept fixed for the different choices of κ and φ .

To check the upper bound in Theorem 3.2, we generate $\mathbf{b}$ and $\tilde{\mathbf{b}}$ using 9 equi-spaced values for φ and κ between 0 and 1, which will provide 81 plots for each sketching strategy. We use a sequence of $L = 10$ independent sketching operators in our BFB approach. After computing values of ν for every case, we evaluate the optimality coefficient difference $\mu(\mathbf{b}, \mathbf{S}_{\ell^*}) - \mu(\mathbf{b}, \mathbf{S}_{\ell^{**}})$. Figure 1 illustrates the relation between $\mu(\mathbf{b}, \mathbf{S}_{\ell^*}) - \mu(\mathbf{b}, \mathbf{S}_{\ell^{**}})$ and the bound $2\sqrt{6(1-\nu)}\varepsilon$. Due to the unknown constants in (3.23) and (3.24), an exact value of ε corresponding to $m = 100$ is unavailable. Instead, we choose ε to be 0.01 heuristically. We chose this particular value of ε since it illustrates how the green curve's shape, which is independent with the scalar ε , separates most of the scatter plots from the rest of the area. The result shows our purposed BFB bound in Theorem 3.2 is effective and non-vacuous for both Gaussian and leverage score sketchings. It is noticeable that all the dots out of our proposed bound (green) are leverage score sketch spots (blue). The reason is because we set $m = 100$ for both sketch strategies, while leverage score sketch requires a higher m to satisfy the (ε, δ) pair condition, which leads to a higher deviation in μ with fixed m ; see details in Theorem 3.11.

To further validate our theoretical results in Theorem 3.5, we consider four combinations of κ and φ as listed in Table 1. For both the Gaussian and leverage score sketches we draw 100 sketches randomly. The same set of sketches are used for each pair of the vectors $\mathbf{b}$ and $\tilde{\mathbf{b}}$. Figure 2 shows scatter plots of the squared optimality coefficients for the four different pairs of $\mathbf{b}$ and $\tilde{\mathbf{b}}$ and two different sketch types.

Table 1 provides the estimated correlations between $\mu^2(\mathbf{b}, \mathbf{S})$ and $\mu^2(\tilde{\mathbf{b}}, \mathbf{S})$ for each of the eight setups based on the data points in Figure 2. For both sketches, a small value of κ and a large value of φ together

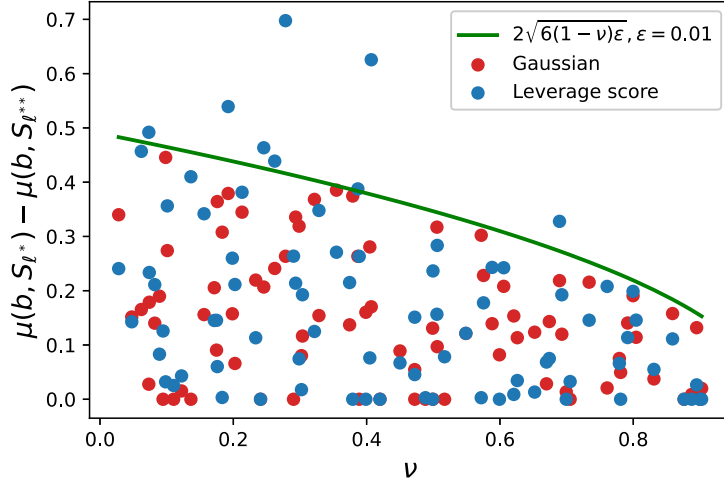


Figure 1: Scatter plots of $\mu(\mathbf{b}, \mathbf{S}_{l^*}) - \mu(\mathbf{b}, \mathbf{S}_{l^{**}})$ based on given values of ν for Gaussian sketch (red) and leverage score sketch (blue). The green curve is the bound we provide in Theorem 3.2 with $\varepsilon = 0.01$.

Table 1: Empirical correlation between $\mu^2(\mathbf{A}, \mathbf{b})$ and $\mu^2(\mathbf{A}, \tilde{\mathbf{b}})$ for four different parameters setups and two different sketch types.

κ	φ	Sketch type	Correlation
0.2	0.3	Gaussian	0.21
0.2	0.95	Gaussian	0.88
0.95	0.3	Gaussian	0.17
0.95	0.95	Gaussian	0.48
0.2	0.3	Leverage score	0.19
0.2	0.95	Leverage score	0.91
0.95	0.3	Leverage score	0.08
0.95	0.95	Leverage score	0.56

yield the highest positive correlation between $\mu^2(\mathbf{b}, \mathbf{S})$ and $\mu^2(\tilde{\mathbf{b}}, \mathbf{S})$. In this case, the sketch that attains the smallest residual on the low-fidelity data also attains a near-minimal residual on the high-fidelity data. This is indicative of the desired sketch transferability between the low- and high-fidelity regression problems. These observations are consistent with the upper bound in (3.3) and the lower bound in (3.13), supporting the idea of BFB.

4.2 Experiments on PDE datasets

In this section we verify the accuracy of Algorithm 2 on two PDE problems: Thermally-driven cavity fluid flow (Section 4.2.1) and simulation of a composite beam (Section 4.2.2). In doing so, we consider three random sketching strategies based on uniform, leverage score (Section 2.2.2), and leveraged volume (Section 2.2.3) sampling. As a baseline, we also present results based on deterministic sketching via column-pivoted QR decomposition (Section 2.2.1).

In both experiments, the high-fidelity solution operator takes uniformly distributed inputs $\mathbf{p} \in [-1, 1]^q$. We therefore consider approximations of the form in (1.1) with $\psi_j : [-1, 1]^q \mapsto \mathbb{R}$ chosen to be products of q univariate (normalized) Legendre polynomials. Specifically, let $\mathbf{j} = (j_1, \dots, j_q)$, $j_k \in \mathbb{N} \cup \{0\}$, be a vector of

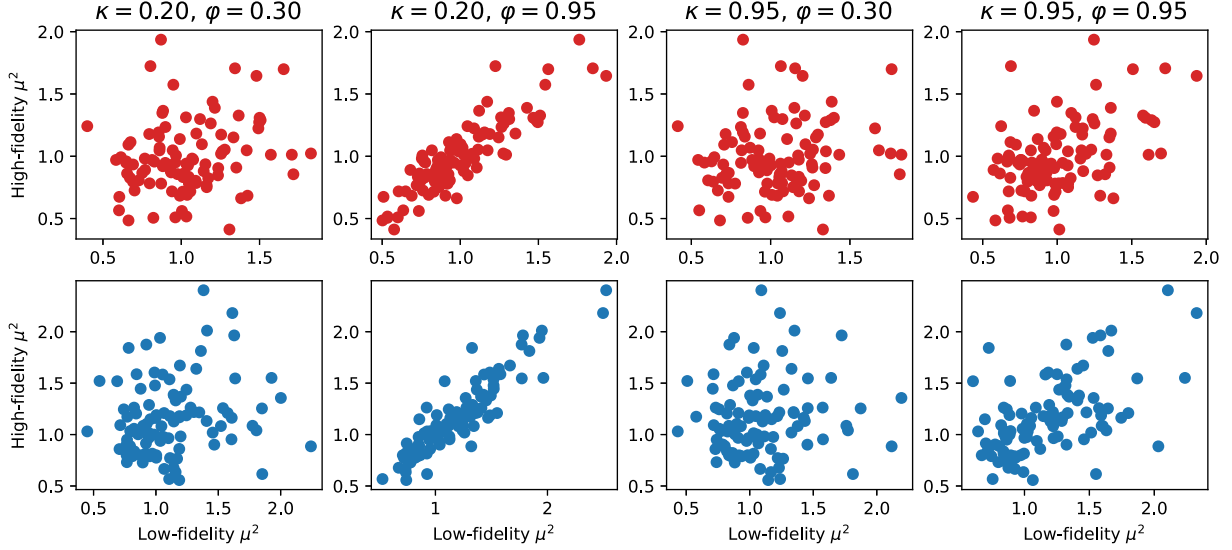


Figure 2: Scatter plots of the square of the optimality coefficient for high- and low-fidelity data for each of 100 different sketches. Each point is equal to $(\mu^2(\tilde{\mathbf{b}}, \mathbf{S}), \mu^2(\mathbf{b}, \mathbf{S}))$ for one realization of the sketch $\mathbf{S}$. The top and bottom panels correspond to the sketches constructed using Gaussian and leverage score sampling sketches, respectively.

non-negative indices and $\psi_{j_k}(p_k)$ denote the Legendre polynomial of degree j_k in p_k such that $\mathbb{E}[\psi_{j_k}^2(p_k)] = 1$. The multivariate Legendre polynomials are given by

$$\psi_{\mathbf{j}}(\mathbf{p}) = \prod_{k=1}^q \psi_{j_k}(p_k).$$

The set of polynomials $\{\psi_{\mathbf{j}}\}$ is chosen so that it spans either a total degree or hyperbolic cross space. In the former case this means all polynomials satisfying $\sum_{k=1}^q j_k \leq \zeta$, while in the latter case $\mathbf{j}$ is limited to multi-indices with $\prod_{k=1}^q (j_k + 1) \leq \zeta + 1$, for some predefined $\zeta \in \mathbb{N} \cup \{0\}$.

In order to construct a design matrix $\mathbf{A}$ as in (1.2), and data vectors $\mathbf{b}$ and $\tilde{\mathbf{b}}$ in (1.2) and (2.8), respectively, we also need to choose pairs of quadrature points and weights $(\mathbf{p}_n, w_n)_{n \in [N]}$. While both deterministic and random rules are possible, we here choose these quantities to be deterministic and of the form

$$\begin{aligned} \mathbf{p}_n &= (p_{1,n_1}, p_{2,n_2}, \dots, p_{q,n_q}), \\ w_n &= \prod_{k=1}^q w_{k,n_k}, \end{aligned} \tag{4.1}$$

where each sequence $(p_{k,n_k}, w_{k,n_k})_{n_k \in [N_k]}$ consists of node-weight pairs in the N_k -point Gauss–Legendre quadrature on $[-1, 1]$. The resulting sequence $(\mathbf{p}_n, w_n)_{n \in [N]}$ contains $N = \prod_{k=1}^q N_k$ pairs. When $\mathbf{A}$ is constructed in this fashion, it is possible to sample rows of that matrix according to the exact leverage score using the efficient method by [Mal+22]. Please see Appendix A for details on how this is done.

To measure the final performance, we use the relative error defined as

$$E \stackrel{\text{def}}{=} \frac{\|\mathbf{A}\hat{\mathbf{x}}_{\text{BFB}} - \mathbf{b}\|_2}{\|\mathbf{b}\|_2},$$

where $\hat{\mathbf{x}}_{\text{BFB}}$ is the output from Algorithm 2.

4.2.1 Cavity fluid flow

Here we consider the case of temperature-driven fluid flow in a 2D cavity [BC15; PHD14; HD15b; HD15a; HD18b], with the quantity of interest being the heat flux averaged along the hot wall as Figure 3 shows. The wall on the left hand side is the hot wall with random temperature T_h , and the cold wall at the right hand side has temperature $T_c < T_h$. $\bar{T}_c$ is the constant mean of T_c . The horizontal walls are adiabatic. The reference temperature and the temperature difference are given by $T_{\text{ref}} = (T_h + \bar{T}_c)/2$ and $\Delta T_{\text{ref}} = T_h - \bar{T}_c$, respectively. The normalized governing equations are given by

$$\begin{aligned} \frac{\partial \mathbf{u}}{\partial t} + \mathbf{u} \cdot \nabla \mathbf{u} &= -\nabla p + \frac{\text{Pr}}{\sqrt{\text{Ra}}} \nabla^2 \mathbf{u} + \text{Pr} \Theta \mathbf{e}_y, \\ \nabla \cdot \mathbf{u} &= 0, \\ \frac{\partial \Theta}{\partial t} + \nabla \cdot (\mathbf{u} \Theta) &= \frac{1}{\sqrt{\text{Ra}}} \nabla^2 \Theta, \end{aligned} \tag{4.2}$$

where $\mathbf{e}_y$ is the unit vector $(0, 1)$, $\mathbf{u} = (u, v)$ is the velocity vector field, $\Theta = (T - T_{\text{ref}})/\Delta T_{\text{ref}}$ is normalized temperature, p is pressure, and t is time. We assume no-slip boundary conditions on the walls. The dimensionless Prandtl and Rayleigh numbers are defined as $\text{Pr} = \nu_{\text{visc}}/\alpha$ and $\text{Ra} = g\tau\Delta T_{\text{ref}}W^3/(\nu_{\text{visc}}\alpha)$, respectively, where W is the width of the cavity, g is gravitational acceleration, ν_{visc} is kinematic viscosity, α is thermal diffusivity, and τ is the coefficient of thermal expansion. We set $g = 10$, $W = 1$, $\tau = 0.5$, $\Delta T_{\text{ref}} = 100$, $\text{Ra} = 10^6$, and $\text{Pr} = 0.71$. On the cold wall, we apply a temperature distribution with stochastic fluctuations as

$$T(x = 1, y) = \bar{T}_c + \sigma_T \sum_{i=1}^q \sqrt{\lambda_i} \phi_i(y) \mu_i,$$

where $\bar{T}_c = 100$ is a constant, $\{\lambda_i\}_{i \in [q]}$ and $\{\phi_i(y)\}_{i \in [q]}$ are the q largest eigenvalues and corresponding eigenfunctions of the kernel $k(y_1, y_2) = \exp(-|y_1 - y_2|/0.15)$, and each $\mu_i \stackrel{\text{i.i.d.}}{\sim} U[-1, 1]$. We let $q = 2$ (though in general, this does not need to match the physical dimension) and $\sigma_T = 2$. The vector $\mathbf{p} = (\mu_1, \mu_2)$ is the uncertain input of the model.

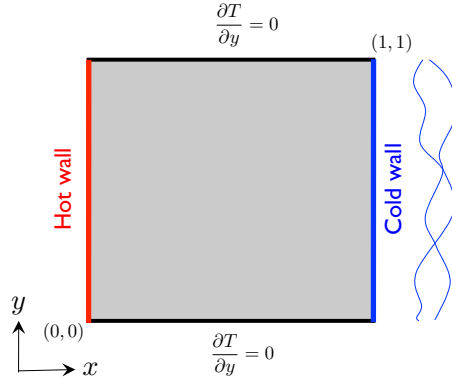


Figure 3: A figure of the temperature driven cavity flow problem, reproduced from Figure 5 of [Fai+17].

In order to solve (4.2) we use the finite volume method with two different grid resolutions: a finer grid of size 128×128 to produce the high-fidelity solution and a coarser grid of size 16×16 to produce the low-fidelity solution. For our surrogate model, we choose the basis set $\{\psi_j\}_{j \in [d]}$ based on the total degree and hyperbolic cross spaces of maximum order $\zeta = 4$. The corresponding spaces have $d = 15$ and $d = 10$ basis functions, respectively. The quadrature pairs $(\mathbf{p}_n, w_n)$ used to construct $\mathbf{A}$, $\mathbf{b}$, and $\tilde{\mathbf{b}}$ are defined as in (4.1) and are based on the nodes and weights from a 10-point Gauss-Legendre rule, i.e., $N_1 = N_2 = 10$.

We first repeat the test we ran on synthetic data in Section 4.1. Figure 4 shows the scatter plots of $(\mu^2(\tilde{\mathbf{b}}, \mathbf{S}), \mu^2(\mathbf{b}, \mathbf{S}))$ for the two different polynomial spaces and three different random sampling approaches. Each plot is based on 100 sketches with $m = 30$ and $m = 20$ samples used for the total degree and hyperbolic cross spaces, respectively. Table 2 presents the correlation coefficients between $\mu^2(\mathbf{b}, \mathbf{S})$ and $\mu^2(\tilde{\mathbf{b}}, \mathbf{S})$ based on the points in Figure 4. There is a discrepancy between the correlation observed for the total degree and hyperbolic cross spaces. One possible explanation for this is that a greater portion of $\mathbf{b}$ is in the range of $\mathbf{A}$ for the total degree space than for the hyperbolic cross space, i.e., κ (see (3.12)) is larger for the former space. Theorem 3.5 indicates that a larger κ should be associated with lower correlation.

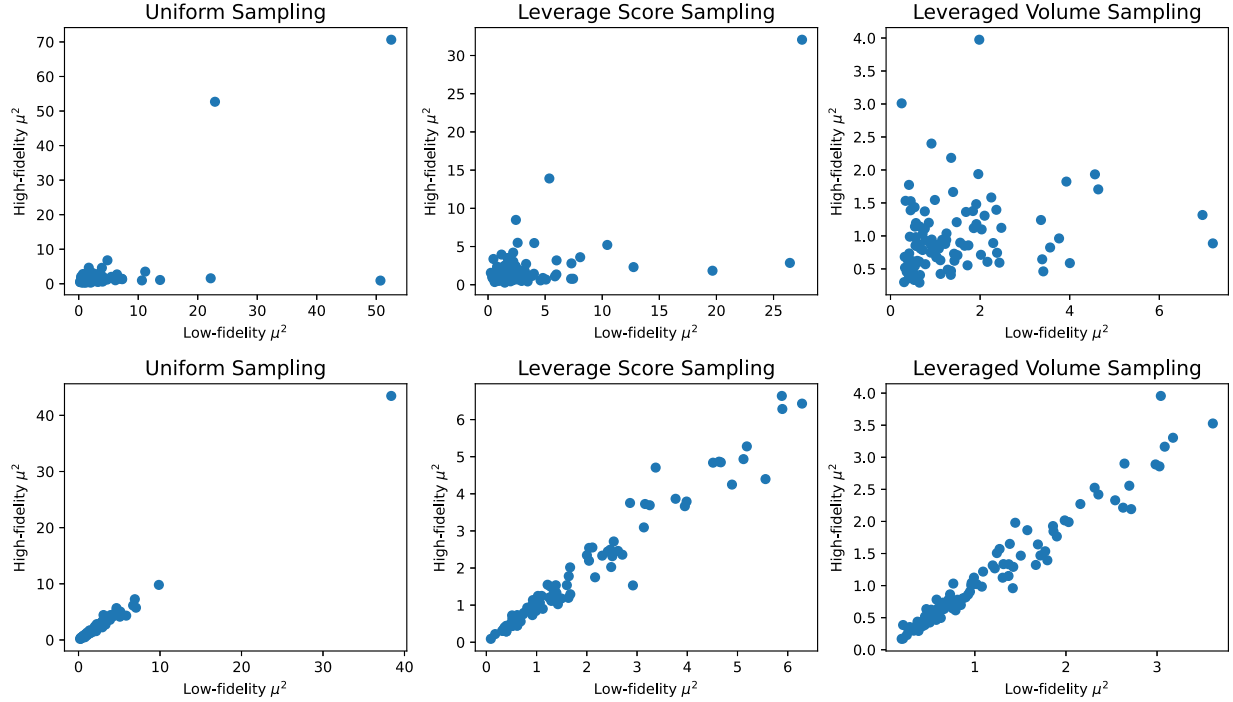


Figure 4: Scatter plots of the square of the optimality coefficient for high- and low-fidelity data from the cavity fluid flow problem for different polynomial spaces (top: total degree; bottom: hyperbolic cross) and types of sampling. Each point is equal to $(\mu^2(\tilde{\mathbf{b}}, \mathbf{S}), \mu^2(\mathbf{b}, \mathbf{S}))$ for one realization of the sketch $\mathbf{S}$, and each subplot contains 100 points (i.e., is based on 100 sketch realizations). For the total degree space $m = 30$ samples are used and for the hyperbolic cross space $m = 20$ samples are used. The corresponding correlation coefficients are presented in Table 2.

Next, we run Algorithm 2 with $L = 10$ sketches and the number of samples $m = 1.2d$ and $m = 2d$. Figure 5 shows the relative error E in (4.2) from running the algorithm 1000 times for each of the different choices of polynomial space, sketch size m , and random sampling approach. We observe that in all cases the BFB approach improves the error as compared to the non-boosted case. In particular, the improvement is more considerable in the case of the hyperbolic cross basis, which is explained by the higher correlation between $\mu^2(\mathbf{A}, \mathbf{b})$ and $\mu^2(\mathbf{A}, \tilde{\mathbf{b}})$, as reported in Table 2. Additionally, for the case of hyperbolic space, the BFB results is comparable or better performance as compared to the column-pivoted QR decomposition (blue line in Figure 5). Note that the computational cost of column-pivoted QR is higher than the BFB as it requires the QR decomposition of the entire matrix $\mathbf{A}$.

Table 2: Correlation coefficients between $\mu^2(\mathbf{A}, \mathbf{b})$ and $\mu^2(\mathbf{A}, \tilde{\mathbf{b}})$ for different sampling methods under total degree or hyperbolic cross space. The correlation is computed based on the points shown in Figure 4.

Polynomial Space	Uniform Sampling	Leverage Score Sampling	Leveraged Volume Sampling
Total Degree	0.66	0.57	0.18
Hyperbolic Cross	0.99	0.98	0.98

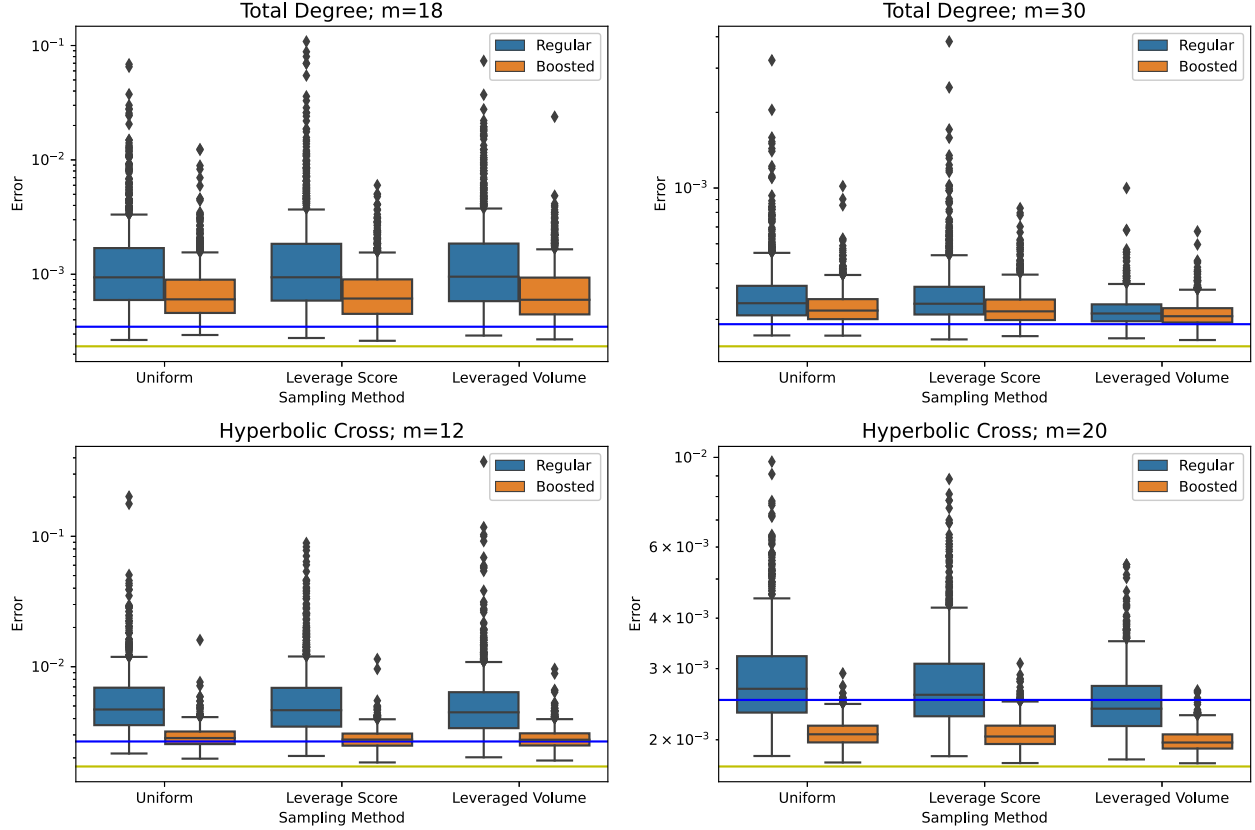


Figure 5: Relative error for different sampling methods and polynomial spaces when fitting the surrogate model to the cavity fluid flow data. Yellow lines show the relative error E in (4.2) for the unsketched solution in (1.2). Blue lines show E when the coefficients $\mathbf{x}$ are computed via the QR decomposition-based method in Section 2.2.1. The blue box plots shows the distribution of E based on 1000 trials when $\mathbf{x}$ is computed as in (2.2). The orange box plots shows the same things, but for the solution $\hat{\mathbf{x}}_{\text{BFB}}$ computed via Algorithm 2.

4.2.2 Composite beam

Following [Ham+18a; De+20; DD22], we consider a plane-stress, cantilever beam with composite cross section and hollow web as shown in Figure 6. The quantity of interest in this case is the maximum displacement of the top cord. The uncertain parameters of the model are E_1, E_2, E_3, f , where E_1, E_2 and E_3

are the Young's moduli of the three components of the cross section and f is the intensity of the applied distributed force on the beam; see Figure 6. These are assumed to be statistically independent and uniformly distributed. The dimension of the input parameter is therefore $q = 4$. Table 3 shows the range of the input parameters as well as the other deterministic parameters.

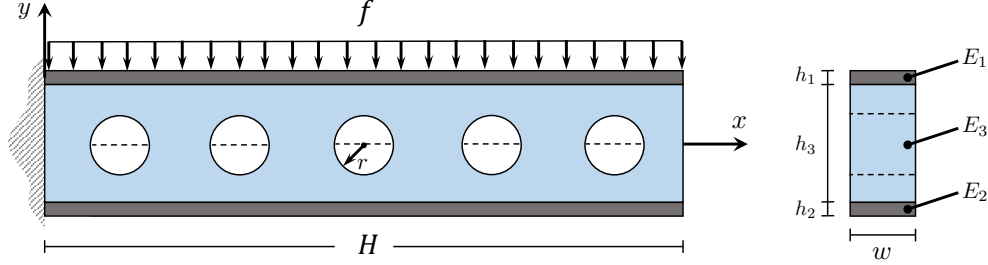


Figure 6: Cantilever beam (left) and the composite cross section (right) adapted from [Ham+18b].

Table 3: The values of the parameters in the composite cantilever beam model. The center of the holes are at $x = \{5, 15, 25, 35, 45\}$. The parameters f , E_1 , E_2 and E_3 are drawn independently and uniformly at random from the specified intervals.

H	h_1	h_2	h_3	w	r	f	E_1	E_2	E_3
50	0.1	0.1	5	1	1.5	[9, 11]	[0.9e6, 1.1e6]	[0.9e6, 1.1e6]	[0.9e4, 1.1e4]

For the cavity fluid flow problem in Section 4.2.1, we created high- and low-fidelity solutions by changing the resolution of the grid used in the numerical solver. For the present problem, we instead use two different models. The high-fidelity model is based on a finite element discretization of the beam using a triangle mesh, as Figure 7 shows. The low-fidelity model is derived from Euler–Bernoulli beam theory in which the vertical cross sections are assumed to remain planes throughout the deformation. The low-fidelity model ignores the shear deformation of the web and does not take the circular holes into account. Considering the Euler-Bernoulli theorem, the vertical displacement u is

$$EI \frac{d^4 u(x)}{dx^4} = -f, \quad (4.3)$$

where E and I are, respectively, the Young's modulus and the moment of inertia of an equivalent cross section consisting of a single material. We let $E = E_3$, and the width of the top and bottom sections are $w_1 = (E_1/E_3)w$ and $w_2 = (E_2/E_3)w$, while all other dimensions are the same, as Figure 6 shows. The solution of (4.3) is

$$u(x) = -\frac{qH^4}{24EI} \left(\left(\frac{x}{H} \right)^4 - 4 \left(\frac{x}{H} \right)^3 + 6 \left(\frac{x}{H} \right)^2 \right).$$



Figure 7: Finite element mesh used to generate high-fidelity solutions.

The surrogate model is based on multivariate Legendre polynomials of maximum degree $\zeta = 2$ with total degree and hyperbolic cross truncation. The corresponding spaces have $d = 15$ and $d = 9$ basis functions, respectively. As in the case of the cavity flow problem, the quadrature pairs $(\mathbf{p}_n, w_n)$ used to construct $\mathbf{A}$, $\mathbf{b}$ and $\tilde{\mathbf{b}}$ are based on the nodes and weights from 10-point Gauss–Legendre rule appropriately mapped into the ranges given in Table 3.

Figure 8 shows the scatter plots of $(\mu^2(\tilde{\mathbf{b}}, \mathbf{S}), \mu^2(\mathbf{b}, \mathbf{S}))$ when repeating the experiment in Section 4.1 for the two different polynomial spaces and three different random sampling approaches. Each plot is based on 100 sketches with $m = 2d$, i.e., $m = 30$ and $m = 18$ samples used for the total degree and hyperbolic cross spaces, respectively. Table 4 reports the correlation coefficient between $\mu^2(\mathbf{b}, \mathbf{S})$ and $\mu^2(\tilde{\mathbf{b}}, \mathbf{S})$, indicating an overall high correlation in all cases.

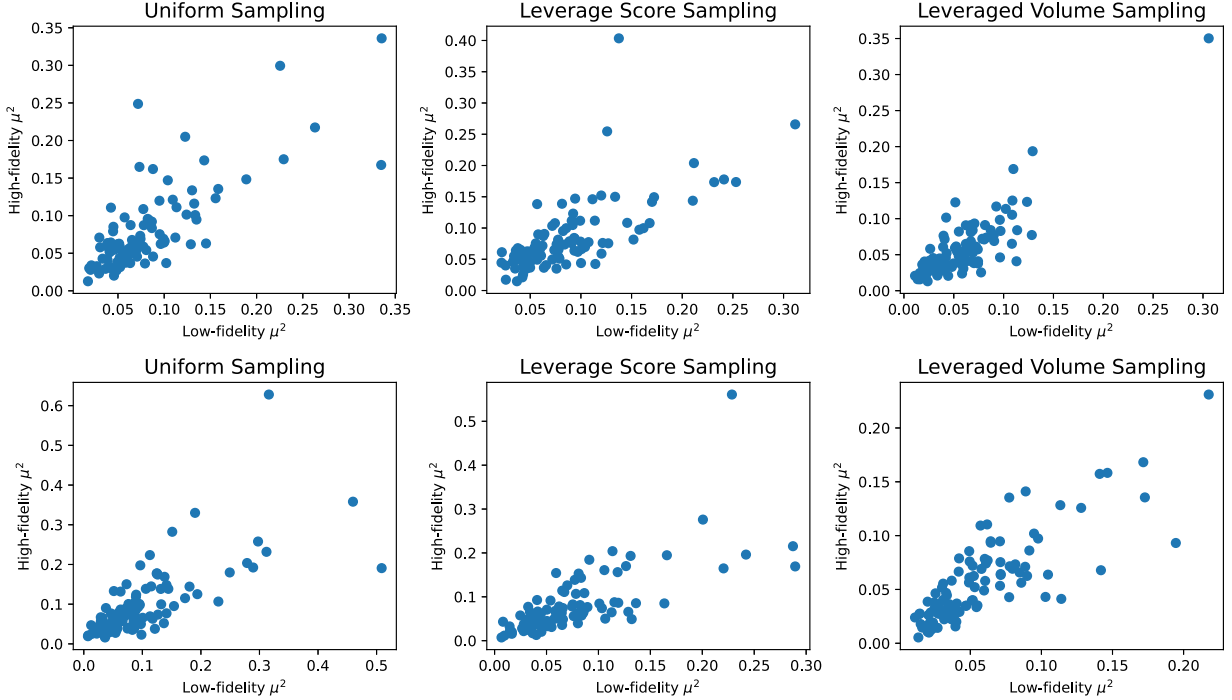


Figure 8: Scatter plots of the square of the optimality coefficient for high- and low-fidelity data from the composite beam problem for different polynomial spaces (top: total degree; bottom: hyperbolic cross) and types of sampling. Each point is equal to $(\mu^2(\tilde{\mathbf{b}}, \mathbf{S}), \mu^2(\mathbf{b}, \mathbf{S}))$ for one realization of the sketch $\mathbf{S}$, and each subplot contains 100 points (i.e., is based on 100 sketch realizations). For the total degree space $m = 30$ samples are used and for the hyperbolic cross space $m = 18$ samples are used. The corresponding correlation coefficients are presented in Table 4.

Table 4: Correlation coefficient between $\mu^2(\mathbf{A}, \mathbf{b})$ and $\mu^2(\mathbf{A}, \tilde{\mathbf{b}})$ for different sampling methods under total degree or hyperbolic cross space. The correlation is computed based on the points shown in Figure 8.

Polynomial Space	Uniform Sampling	Leverage Score Sampling	Leveraged Volume Sampling
Total Degree	0.77	0.69	0.84
Hyperbolic Cross	0.72	0.73	0.82

Next, we run Algorithm 2 with $L = 10$ sketches and m chosen to be $m = 1.2d$ and $m = 2d$. Figure 9 shows the results from running the algorithm 1000 times for each of the different choices of polynomial space, number of samples m , and random sampling approach. We observe that the BFB performance is superior to that of the non-boosted implementation as it leads to smaller variance of the error and fewer outliers with smaller deviation from the mean performance. In this example, the BFB leads to comparable accuracy as the column-pivoted QR sketch, but with smaller sketching cost. As in the case of the cavity flow, the results corroborate the discussion below Theorem 3.5, in that the BFB improves the regression accuracy when $\text{corr}(\mu^2(\mathbf{b}, \mathbf{S}), \mu(\tilde{\mathbf{b}}, \mathbf{S}))$ is large.

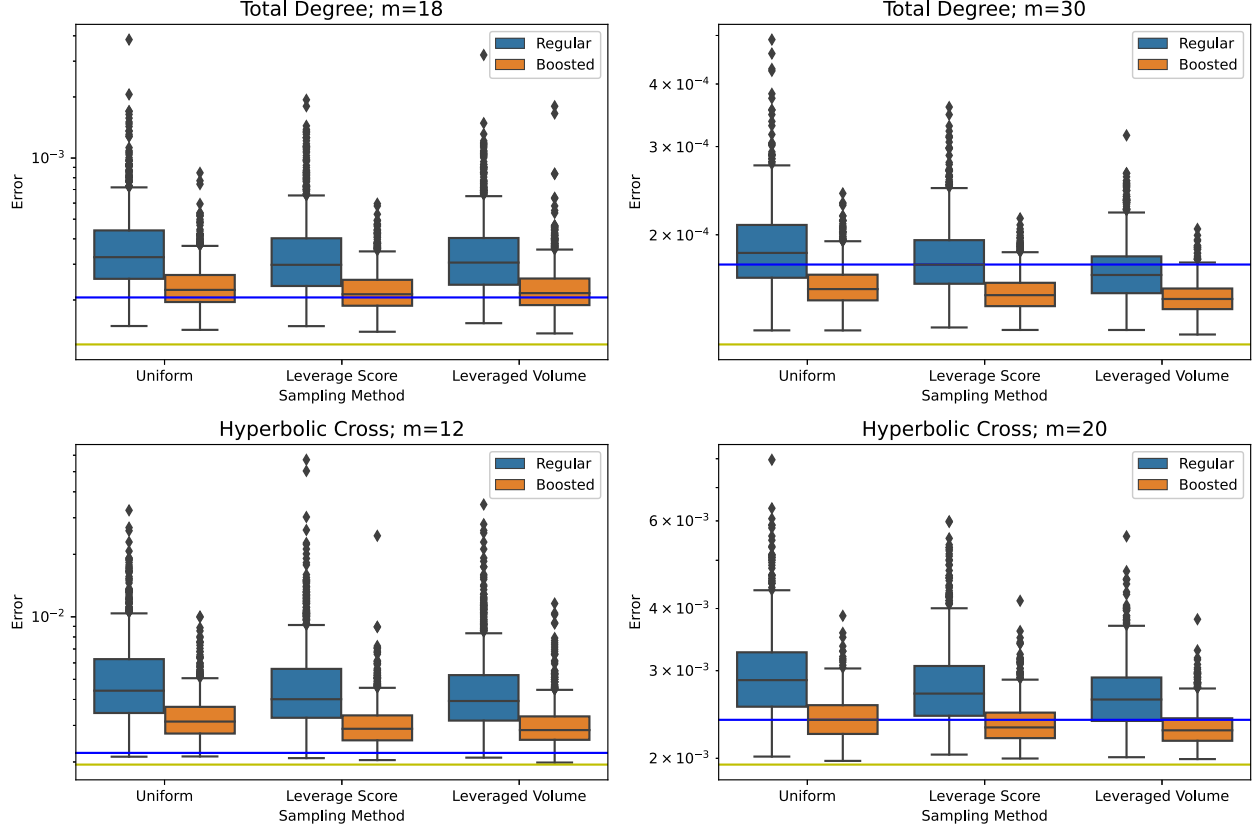


Figure 9: Relative error for different sampling methods and polynomial spaces when fitting the surrogate model to the beam problem data. Yellow lines show the relative error E in (4.2) for the unsketched solution in (1.2). Blue lines show E when the coefficients $\mathbf{x}$ are computed via the QR decomposition-based method in Section 2.2.1. The blue box plots shows the distribution of E based on 1000 trials when $\mathbf{x}$ is computed as in (2.2). The orange box plots shows the same things, but for the solution $\hat{\mathbf{x}}_{\text{BFB}}$ computed via Algorithm 2.

5 Conclusion

This work was concerned with the construction of (polynomial) emulators of parameter-to-solution maps of PDE problems via sketched least-squares regression. Sketching is a design of experiments approach that aims to improve the cost of building a least squares solution in terms of reducing the number of samples needed — when the cost of generating data is high — or the cost of generating a least squares solution — when data size is substantial. Focusing on the former case, we have proposed a new boosting algorithm to compute a sketched least squares solution.

The procedure consisted in identifying the best sketch from a set of candidates used to construct least squares regression of the low-fidelity data and applying this *optimal* sketch to the regression of high-fidelity data. The bi-fidelity boosting (BFB) approach limits the required sample complexity to $\sim d \log d$ high-fidelity data, where d is the size of the (polynomial) basis. We have provided theoretical analysis of the BFB approach identifying assumptions on the low- and high-fidelity data under which the BFB leads to improvement of the solution relative to non-boosted regression of the high-fidelity data. We have also provided quantitative bounds on the residual of the BFB solution relative to the full, computationally expensive solution. We have

investigated the performance of BFB on manufactured and PDE data from fluid and solid mechanics. These cover sketching strategies based on leverage score and leveraged volume sampling, for truncated Legendre polynomials of both total degree and hyperbolic cross type. All tests illustrated the efficacy of BFB in reducing the residual — as compared to the non-boosted implementation — and validate the theoretical results.

The present study was focused on the case of (weighted) least squares polynomial regression. When the regression coefficients are sparse, methods based on compressive sampling have proven efficient in reducing the sample complexity below the size of the polynomial basis; see, e.g., [DO11; ABW22]. As interesting future research direction is to extend the BFB strategy to such under-determined cases, for instance, using the approach of [DDH18].

Acknowledgments

This work was supported by the AFOSR awards FA9550-20-1-0138 and FA9550-20-1-0188 with Dr. Fariba Fahroo as the program manager. The views expressed in the article do not necessarily represent the views of the AFOSR or the U.S. Government.

A Efficient leverage score sampling of certain design matrices

In this section, we describe the key elements of the sampling approach developed in [Mal+22] as it applies to the problems we consider in this paper. The discussion here will consider the design matrices discussed in Section 4.2. Using the same notation as in that section, define the matrices $\mathbf{A}_k$ for $k \in [q]$ elementwise via

$$\mathbf{A}_k(n_k, j_k) = \sqrt{w_{k,n_k}} \psi_{j_k}(p_{k,n_k}), \quad n_k \in [N_k], j_k \in [\zeta].$$

Next, define

$$\mathbf{A}_{\text{TP}} \stackrel{\text{def}}{=} \mathbf{A}_1 \otimes \cdots \otimes \mathbf{A}_q,$$

where $\otimes$ denotes the Kronecker product; see Section 12.3 of [GVL13] for a definition. The design matrices corresponding to total degree and hyperbolic cross polynomial spaces discussed in Section 4.2 are made up of a subset of the columns of $\mathbf{A}_{\text{TP}}$. In particular, using Matlab indexing notation, they can be written as

$$\mathbf{A} = \mathbf{A}_{\text{TP}}(:, \mathbf{v}),$$

where $\mathbf{v}$ is a vector containing distinct column indices of $\mathbf{A}_{\text{TP}}$. The sampling scheme we discuss requires the additional assumption that the entries in $\mathbf{v}$ are arranged in increasing order. The columns of $\mathbf{A}$ can always be permuted to ensure that this is possible when $\mathbf{A}$ corresponds to a total degree or hyperbolic cross space. Such a permutation will not change the least squares problem since it will only permute the order of the entries in the solution vector, and is therefore something that can always be done.

Note that a column index c of $\mathbf{A}_{\text{TP}}$ corresponds to a multi-index $(c_1, \dots, c_q)$ such that

$$\mathbf{A}_{\text{TP}}(:, c) = \mathbf{A}_1(:, c_1) \otimes \cdots \otimes \mathbf{A}_q(:, c_q).$$

Each row index r of $\mathbf{A}_{\text{TP}}$ corresponds to a multi-index $(r_1, \dots, r_q)$ in a similar fashion.

Algorithm 3 outlines the sampling algorithm. We provide some intuition for why the algorithm works and refer the reader to [Mal+22] for a rigorous treatment. Note that $\mathbf{A}$ is full rank and therefore $\text{rank}(\mathbf{A}) = d$. Let $\mathbf{Q}\mathbf{R} = \mathbf{A}$ be a compact QR decomposition (i.e., such that $\mathbf{Q}$ has d columns and $\mathbf{R}$ has d rows). Recall that the leverage score sampling distribution satisfies

$$\mathbf{p}(i) = \frac{\|\mathbf{Q}(i, :)\|_2^2}{d}.$$

Instead of drawing a sample according to the distribution above, we may instead draw a single column $\mathbf{Q}(:, j)$ of $\mathbf{Q}$ uniformly at random and instead draw a sample according to the probability distribution defined by $\tilde{\mathbf{p}}_j(i) = (\mathbf{Q}(i, j))^2$. To see this, let $\tilde{I}$ be a random row index drawn according to this alternate strategy. Moreover, let $J \sim \text{Uniform}([d])$ be the random column index, and let $\tilde{I}_j$ be a random row index drawn according to $\tilde{\mathbf{p}}_j$. Then we have

$$\mathbb{P}(\tilde{I} = i) = \sum_{j=1}^d \mathbb{P}(\tilde{I} = i \mid J = j) \mathbb{P}(J = j) = \sum_{j=1}^d \mathbb{P}(\tilde{I}_j = i) \mathbb{P}(J = j) = \sum_{j=1}^d (\mathbf{Q}(i, j))^2 \frac{1}{d} = \frac{\|\mathbf{Q}(i, :)\|_2^2}{d} = \mathbf{p}(i).$$

This shows that the alternate sampling strategy indeed draws samples according to the leverage score sampling distribution. This is the sampling strategy that our algorithm uses. Moreover, it uses two additional fact:

- (i) When $\mathbf{A}$ has the particular structure assumed in this section, then the c th column of $\mathbf{Q}$ satisfies

$$\mathbf{Q}(:, c) = \mathbf{Q}_1(:, c_1) \otimes \cdots \otimes \mathbf{Q}_q(:, c_q), \tag{A.1}$$

where $\mathbf{Q}_1, \dots, \mathbf{Q}_q$ are defined in line 2 in Algorithm 3.

- (ii) Due to (A.1), drawing a row index r according to $\tilde{\mathbf{p}}_j$ is equivalent to drawing a multi-index $(r_1, \dots, r_q)$ according to a product distribution with each r_k drawn independently according to the distribution $((\mathbf{Q}_k(r_k, j_k))^2)_{r_k}$ where $(j_1, \dots, j_q)$ is the column multi-index corresponding to j .

Fact (i) makes it possible to sample according to the alternate sampling strategy without every needing to compute the QR decomposition of the large matrix $\mathbf{A}$. A more general version of this fact appears in Proposition 4.4 of [Mal+22]. Fact (ii) further makes it possible to sample according to $\tilde{\mathbf{p}}_j$ without needing to form that probability vector which is of length $\prod_k N_k$.

Algorithm 3: Efficient leverage score sampling of total degree and hyperbolic cross design matrices

Input: Matrices $\mathbf{A}_1, \dots, \mathbf{A}_q$, index vector $\mathbf{v}$, number of samples m

Output: Vector $\mathbf{s} \in [\prod_k N_k]^m$ of m samples drawn from row indices of $\mathbf{A}$

```

1: for  $k \in [q]$  do
2:   Compute compact QR decomposition  $\mathbf{Q}_k \mathbf{R}_k = \mathbf{A}_k$ 
3: end for
4: for  $i \in [m]$  do
5:   Draw an entry  $j$  from  $\mathbf{v}$  uniformly at random
6:   Compute the multi-index  $(j_1, \dots, j_q)$  corresponding to  $j$ 
7:   for  $k \in [q]$  do
8:     Construct the probability distribution  $\mathbf{p} = ((\mathbf{Q}_k(r_k, j_k))^2)_{r_k} \in \mathbb{R}^{N_k}$ 
9:     Draw an index  $r_k \in [N_k]$  according to the distribution  $\mathbf{p}$ 
10:  end for
11:  Set the  $i$ th sample  $\mathbf{s}(i)$  equal the row index corresponding to the row multi-index  $(r_1, \dots, r_q)$ 
12: end for
13: return Vector of samples  $\mathbf{s}$ 

```

B Proof of Theorem 3.5

The proof of Theorem 3.5 relies on the following lemmas:

Lemma B.1. *Let X and Y be two (nonconstant) random variables defined on the same probability space. The correlation coefficient between X and Y , $\text{corr}(X, Y)$, is bounded from below as*

$$\text{corr}(X, Y) \geq \sqrt{\frac{\mathbb{V}[X]}{\mathbb{V}[Y]}} - \sqrt{\frac{\mathbb{V}[Y - X]}{\mathbb{V}[Y]}}.$$

Proof. It follows from direct computation that

$$\begin{aligned}
\text{corr}(X, Y) &= \frac{\mathbb{E}[XY] - \mathbb{E}[X]\mathbb{E}[Y]}{\sqrt{\mathbb{V}[X]\mathbb{V}[Y]}} \\
&= \frac{\mathbb{E}[X^2] - \mathbb{E}[X]^2}{\sqrt{\mathbb{V}[X]\mathbb{V}[Y]}} + \frac{\mathbb{E}[X(Y - X)] - \mathbb{E}[X]\mathbb{E}[Y - X]}{\sqrt{\mathbb{V}[X]\mathbb{V}[Y]}} \\
&= \sqrt{\frac{\mathbb{V}[X]}{\mathbb{V}[Y]}} + \text{corr}(X, Y - X) \sqrt{\frac{\mathbb{V}[Y - X]}{\mathbb{V}[Y]}} \\
&\geq \sqrt{\frac{\mathbb{V}[X]}{\mathbb{V}[Y]}} - \sqrt{\frac{\mathbb{V}[Y - X]}{\mathbb{V}[Y]}},
\end{aligned}$$

where the last inequality uses $\text{corr}(X, Y - X) \geq -1$. □

Lemma B.2. *Let $\boldsymbol{\xi} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_n)$ be a standard Gaussian vector in $\mathbb{R}^n$. For any $\mathbf{w}, \mathbf{z} \in \mathbb{R}^n$,*

$$\mathbb{E}[\langle \mathbf{w}, \boldsymbol{\xi} \rangle^2 \langle \mathbf{z}, \boldsymbol{\xi} \rangle^2] = 2\langle \mathbf{w}, \mathbf{z} \rangle^2 + \|\mathbf{w}\|_2^2 \|\mathbf{z}\|_2^2. \quad (\text{B.1})$$

Proof. The proof follows from a direct application of Wick's formula [Wic50]. Denote $X_1 = \langle \mathbf{w}, \boldsymbol{\xi} \rangle$ and $X_2 = \langle \mathbf{z}, \boldsymbol{\xi} \rangle$. It is easy to verify that

$$\begin{pmatrix} X_1 \\ X_2 \end{pmatrix} \sim \mathcal{N}(\mathbf{0}, \mathbf{K}), \quad \mathbf{K} = \begin{pmatrix} \|\mathbf{w}\|_2^2 & \langle \mathbf{w}, \mathbf{z} \rangle \\ \langle \mathbf{w}, \mathbf{z} \rangle & \|\mathbf{z}\|_2^2 \end{pmatrix}.$$

By Wick's formula,

$$\mathbb{E}[\langle \mathbf{w}, \boldsymbol{\xi} \rangle^2 \langle \mathbf{z}, \boldsymbol{\xi} \rangle^2] = \mathbb{E}[X_1^2 X_2^2] = 2\mathbb{E}[X_1 X_2]^2 + \mathbb{E}[X_1^2] \mathbb{E}[X_2^2] = 2\langle \mathbf{w}, \mathbf{z} \rangle^2 + \|\mathbf{w}\|_2^2 \|\mathbf{z}\|_2^2.$$

□

Proof of Theorem 3.5. Since correlation coefficients are scale-invariant, and both $\mathbf{b}$ and $\tilde{\mathbf{b}}$ are fixed,

$$\text{corr}(\mu^2(\mathbf{b}, \mathbf{S}), \mu^2(\tilde{\mathbf{b}}, \mathbf{S})) = \text{corr}\left(\frac{r_{\mathbf{S}}^2(\mathbf{A}, \mathbf{b}) - r^2(\mathbf{A}, \mathbf{b})}{\|\mathbf{b}\|_2^2}, \frac{r_{\mathbf{S}}^2(\mathbf{A}, \tilde{\mathbf{b}}) - r^2(\mathbf{A}, \tilde{\mathbf{b}})}{\|\tilde{\mathbf{b}}\|_2^2}\right).$$

Without loss of generality, we assume $\|\mathbf{b}\|_2 = \|\tilde{\mathbf{b}}\|_2 = 1$, so that $\mathbf{b}_{\mathcal{P}} = \mathbf{b}$, $\tilde{\mathbf{b}}_{\mathcal{P}} = \tilde{\mathbf{b}}$.

Let

$$\begin{aligned} X &= r_{\mathbf{S}}^2(\mathbf{A}, \mathbf{b}) - r^2(\mathbf{A}, \mathbf{b}) = \|(\mathbf{S}\mathbf{Q})^\dagger \mathbf{S}\mathbf{Q}_\perp \mathbf{Q}_\perp^T \mathbf{b}\|_2^2 \\ Y &= r_{\mathbf{S}}^2(\mathbf{A}, \tilde{\mathbf{b}}) - r^2(\mathbf{A}, \tilde{\mathbf{b}}) = \|(\mathbf{S}\mathbf{Q})^\dagger \mathbf{S}\mathbf{Q}_\perp \mathbf{Q}_\perp^T \tilde{\mathbf{b}}\|_2^2. \end{aligned}$$

To apply Lemma B.1, it suffices to estimate $\mathbb{V}[X]/\mathbb{V}[Y]$ and $\mathbb{V}[Y - X]/\mathbb{V}[Y]$.

First of all, due to the rotation invariance of joint Gaussians,

$$\begin{aligned} \mathbf{G}_1 &\stackrel{\text{def}}{=} \sqrt{m} \mathbf{S}\mathbf{Q} \in \mathbb{R}^{m \times d}, \\ \mathbf{G}_2 &\stackrel{\text{def}}{=} \sqrt{m} \mathbf{S}\mathbf{Q}_\perp \in \mathbb{R}^{m \times (N-d)} \end{aligned}$$

are independent Gaussian random matrices, i.e., $(\mathbf{S}\mathbf{Q})^\dagger \mathbf{S}\mathbf{Q}_\perp \mathbf{Q}_\perp^T = \mathbf{G}_1^\dagger \mathbf{G}_2 \mathbf{Q}_\perp^T$, and

$$\begin{aligned} \mathbb{E}[X] &= \mathbb{E}\left[\text{tr}\left(\mathbf{G}_1^\dagger \mathbf{G}_2 \mathbf{Q}_\perp^T \mathbf{b} \mathbf{b}^T \mathbf{Q}_\perp \mathbf{G}_2^T \mathbf{G}_1^\dagger\right)\right] \\ &= \mathbb{E}\left[\text{tr}\left(\mathbf{G}_1^\dagger \mathbb{E}[\mathbf{G}_2 \mathbf{Q}_\perp^T \mathbf{b} \mathbf{b}^T \mathbf{Q}_\perp \mathbf{G}_2^T] \mathbf{G}_1^\dagger\right)\right] \\ &= \|\mathbf{Q}_\perp^T \mathbf{b}\|_2^2 \mathbb{E}\left[\text{tr}\left(\mathbf{G}_1^\dagger \mathbf{G}_1^\dagger\right)\right] \\ &= \|\mathbf{Q}_\perp^T \mathbf{b}\|_2^2 \mathbb{E}\left[\text{tr}((\mathbf{G}_1^T \mathbf{G}_1)^{-1})\right], \end{aligned}$$

where we have used that $\mathbb{E}[\mathbf{G}_2 \mathbf{Q}_\perp^T \mathbf{b} \mathbf{b}^T \mathbf{Q}_\perp \mathbf{G}_2^T] = \|\mathbf{Q}_\perp^T \mathbf{b}\|_2^2 \mathbf{I}_m$.

Note $\mathbf{G}_1^\dagger \mathbf{G}_1$ is a Wishart matrix with dimension d and degrees of freedom m , i.e. $\mathbf{W} = \mathbf{G}_1^T \mathbf{G}_1 \sim W_d(\mathbf{I}_d, m)$. Consequently, $\mathbb{E}[\mathbf{W}^{-1}] = \frac{1}{m-d-1} \mathbf{I}_d$ if $m > d+1$, and

$$\mathbb{E}[X] = \|\mathbf{Q}_\perp^T \mathbf{b}\|_2^2 \frac{d}{m-d-1}. \quad (\text{B.2})$$

Similarly,

$$\mathbb{E}[Y] = \|\mathbf{Q}_\perp^T \tilde{\mathbf{b}}\|_2^2 \frac{d}{m-d-1}. \quad (\text{B.3})$$

Note $\mathbf{G}_1^\dagger \mathbf{G}_2 \mathbf{Q}_\perp^T \mathbf{a} \stackrel{\mathcal{D}}{=} \|\mathbf{Q}_\perp^T \mathbf{a}\|_2 (\mathbf{G}_1^T \mathbf{G}_1)^{-1} \mathbf{G}_1^T \boldsymbol{\xi}$ for every $\mathbf{a} \in \mathbb{R}^N$, where $\boldsymbol{\xi} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_m)$ is independent of $\mathbf{G}_1$. If

we denote $\mathbf{G} = (\mathbf{G}_1^T \mathbf{G}_1)^{-1} \mathbf{G}_1^T$, with rows denoted by $\mathbf{g}_i, i \in [d]$, then

$$\begin{aligned}
\mathbb{E}[X^2] &= \mathbb{E}[\|\mathbf{Q}_\perp^T \mathbf{b}\|_2 \mathbf{G} \boldsymbol{\xi} \|_2^4] \\
&= \|\mathbf{Q}_\perp^T \mathbf{b}\|_2^4 \mathbb{E} \left[\left(\sum_{i=1}^d \langle \mathbf{g}_i, \boldsymbol{\xi} \rangle^2 \right)^2 \right] \\
&= \|\mathbf{Q}_\perp^T \mathbf{b}\|_2^4 \left(\sum_{i=1}^d \mathbb{E}[\langle \mathbf{g}_i, \boldsymbol{\xi} \rangle^4] + \sum_{i \neq j} \mathbb{E}[\langle \mathbf{g}_i, \boldsymbol{\xi} \rangle^2 \langle \mathbf{g}_j, \boldsymbol{\xi} \rangle^2] \right) \\
&\stackrel{(\text{B.1})}{=} \|\mathbf{Q}_\perp^T \mathbf{b}\|_2^4 \left(3 \sum_{i=1}^d \mathbb{E}[\|\mathbf{g}_i\|_2^4] + \sum_{i \neq j} (2\mathbb{E}[\langle \mathbf{g}_i, \mathbf{g}_j \rangle^2] + \mathbb{E}[\|\mathbf{g}_i\|_2^2 \|\mathbf{g}_j\|_2^2]) \right) \\
&= \|\mathbf{Q}_\perp^T \mathbf{b}\|_2^4 \left(2\mathbb{E}[\|\mathbf{G} \mathbf{G}^T\|_F^2] + \mathbb{E}[\text{tr}(\mathbf{G} \mathbf{G}^T)^2] \right) \\
&= \|\mathbf{Q}_\perp^T \mathbf{b}\|_2^4 \left(2\mathbb{E}[\|(\mathbf{G}_1^T \mathbf{G}_1)^{-1}\|_F^2] + \mathbb{E}[\text{tr}((\mathbf{G}_1^T \mathbf{G}_1)^{-1})^2] \right).
\end{aligned} \tag{B.4}$$

To explicitly compute (B.4), we use the following moments formulas of inverse Wishart distributions [Kol05, Theorem 2.4.14]:

$$\begin{aligned}
\mathbb{E}[\mathbf{W}^{-1} \mathbf{W}^{-1}] &= \left(\frac{d}{(m-d)(m-d-3)} + \frac{d}{(m-d)(m-d-1)(m-d-3)} \right) \mathbf{I}_d \\
\text{Cov}(\mathbf{W}_{ii}^{-1}, \mathbf{W}_{jj}^{-1}) &= \frac{2 + 2(m-d-1)\delta_{ij}}{(m-d)(m-d-1)^2(m-d-3)}.
\end{aligned}$$

Therefore,

$$\mathbb{E}[\|(\mathbf{G}_1^T \mathbf{G}_1)^{-1}\|_F^2] = \text{tr}(\mathbb{E}[\mathbf{W}^{-1} \mathbf{W}^{-1}]) = \frac{d^2}{(m-d-1)(m-d-3)} \simeq \frac{d^2}{(m-d-1)^2}$$

and

$$\begin{aligned}
\mathbb{E}[\text{tr}((\mathbf{G}_1^T \mathbf{G}_1)^{-1})^2] &= \sum_{i,j \in [d]} \mathbb{E}[\mathbf{W}_{ii}^{-1} \mathbf{W}_{jj}^{-1}] \\
&= \sum_{i,j \in [d]} (\text{Cov}(\mathbf{W}_{ii}^{-1}, \mathbf{W}_{jj}^{-1}) + \mathbb{E}[\mathbf{W}_{ii}^{-1}] \mathbb{E}[\mathbf{W}_{jj}^{-1}]) \\
&= \frac{d^2}{(m-d-1)^2} + \frac{2d}{(m-d-1)^2(m-d-3)} + \frac{2(d^2-d)}{(m-d)(m-d-1)^2(m-d-3)} \\
&\simeq \frac{d^2}{(m-d-1)^2},
\end{aligned}$$

where $a_m \simeq b_m$ if $\lim_{m \rightarrow \infty} a_m/b_m = 1$. Substituting these back into (B.4) yields

$$\mathbb{E}[X^2] \simeq \|\mathbf{Q}_\perp^T \mathbf{b}\|_2^4 \frac{3d^2}{(m-d-1)^2}. \tag{B.5}$$

Replacing $\mathbf{b}$ by $\tilde{\mathbf{b}}$ in the above computation gives a similar estimate for $\mathbb{E}[Y^2]$:

$$\mathbb{E}[Y^2] \simeq \|\mathbf{Q}_\perp^T \tilde{\mathbf{b}}\|_2^4 \frac{3d^2}{(m-d-1)^2}. \tag{B.6}$$

Combining (B.5), (B.6) with (B.2) and (B.3) produces

$$\begin{aligned}\mathbb{V}[X] &\simeq \|\mathbf{Q}_\perp^T \mathbf{b}\|_2^4 \frac{2d^2}{(m-d-1)^2}, \\ \mathbb{V}[Y] &\simeq \|\mathbf{Q}_\perp^T \tilde{\mathbf{b}}\|_2^4 \frac{2d^2}{(m-d-1)^2},\end{aligned}$$

which implies

$$\frac{\mathbb{V}[X]}{\mathbb{V}[Y]} \simeq \frac{\|\mathbf{Q}_\perp^T \mathbf{b}\|_2^4}{\|\mathbf{Q}_\perp^T \tilde{\mathbf{b}}\|_2^4}.$$

On the other hand, using Cauchy–Schwarz inequality, Moreover, we have

$$\begin{aligned}\mathbb{V}[Y - X] &\leq \mathbb{E}[(Y - X)^2] \\ &= \mathbb{E}[(X^{\frac{1}{2}} + Y^{\frac{1}{2}})^2 \cdot (X^{\frac{1}{2}} - Y^{\frac{1}{2}})^2] \\ &= \mathbb{E}[(X^{\frac{1}{2}} + Y^{\frac{1}{2}})^2 \cdot (\|\mathbf{G}_1^\dagger \mathbf{G}_2 \mathbf{Q}_\perp^T \mathbf{b}\|_2 - \|\mathbf{G}_1^\dagger \mathbf{G}_2 \mathbf{Q}_\perp^T \tilde{\mathbf{b}}\|_2)^2] \\ &\leq \mathbb{E}[(X^{\frac{1}{2}} + Y^{\frac{1}{2}})^2 \cdot \|\mathbf{G}_1^\dagger \mathbf{G}_2 \mathbf{Q}_\perp^T (\mathbf{b} \pm \tilde{\mathbf{b}})\|_2^2],\end{aligned}\tag{B.7}$$

where the last inequality follows from the reverse triangle inequality. Furthermore, using the inequality of arithmetic and geometric means followed by the Cauchy–Schwarz inequality, we have

$$\begin{aligned}&\mathbb{E}[(X^{\frac{1}{2}} + Y^{\frac{1}{2}})^2 \cdot \|\mathbf{G}_1^\dagger \mathbf{G}_2 \mathbf{Q}_\perp^T (\mathbf{b} \pm \tilde{\mathbf{b}})\|_2^2] \\ &\leq 2\mathbb{E}[(X + Y) \cdot \|\mathbf{G}_1^\dagger \mathbf{G}_2 \mathbf{Q}_\perp^T (\mathbf{b} \pm \tilde{\mathbf{b}})\|_2^2] \\ &= 2\mathbb{E}[X \cdot \|\mathbf{G}_1^\dagger \mathbf{G}_2 \mathbf{Q}_\perp^T (\mathbf{b} \pm \tilde{\mathbf{b}})\|_2^2] + 2\mathbb{E}[Y \cdot \|\mathbf{G}_1^\dagger \mathbf{G}_2 \mathbf{Q}_\perp^T (\mathbf{b} \pm \tilde{\mathbf{b}})\|_2^2] \\ &\leq 2\sqrt{\mathbb{E}[X^2] \cdot \mathbb{E}[\|\mathbf{G}_1^\dagger \mathbf{G}_2 \mathbf{Q}_\perp^T (\mathbf{b} \pm \tilde{\mathbf{b}})\|_2^4]} + 2\sqrt{\mathbb{E}[Y^2] \cdot \mathbb{E}[\|\mathbf{G}_1^\dagger \mathbf{G}_2 \mathbf{Q}_\perp^T (\mathbf{b} \pm \tilde{\mathbf{b}})\|_2^4]}.\end{aligned}\tag{B.8}$$

Combining (B.7) and (B.8) yields

$$\mathbb{V}[Y - X] \leq 2\sqrt{\mathbb{E}[X^2] \cdot \mathbb{E}[\|\mathbf{G}_1^\dagger \mathbf{G}_2 \mathbf{Q}_\perp^T (\mathbf{b} \pm \tilde{\mathbf{b}})\|_2^4]} + 2\sqrt{\mathbb{E}[Y^2] \cdot \mathbb{E}[\|\mathbf{G}_1^\dagger \mathbf{G}_2 \mathbf{Q}_\perp^T (\mathbf{b} \pm \tilde{\mathbf{b}})\|_2^4]}.\tag{B.9}$$

A similar argument as (B.5) shows that

$$\mathbb{E}[\|\mathbf{G}_1^\dagger \mathbf{G}_2 \mathbf{Q}_\perp^T (\mathbf{b} \pm \tilde{\mathbf{b}})\|_2^4] \simeq \|\mathbf{Q}_\perp^T (\mathbf{b} \pm \tilde{\mathbf{b}})\|_2^4 \frac{3d^2}{(m-d-1)^2}.\tag{B.10}$$

Plugging (B.10) into (B.9) together with the previous estimates yields that, asymptotically,

$$\frac{\mathbb{V}[Y - X]}{\mathbb{V}[Y]} \leq \frac{2(\|\mathbf{Q}_\perp^T \mathbf{b}\|_2^2 + \|\mathbf{Q}_\perp^T \tilde{\mathbf{b}}\|_2^2) \|\mathbf{Q}_\perp^T (\mathbf{b} \pm \tilde{\mathbf{b}})\|_2^2}{\|\mathbf{Q}_\perp^T \tilde{\mathbf{b}}\|_2^4} \cdot \frac{3d^2}{2d^2} \leq \frac{6\|\mathbf{Q}_\perp^T (\mathbf{b} \pm \tilde{\mathbf{b}})\|_2^2}{\|\mathbf{Q}_\perp^T \tilde{\mathbf{b}}\|_2^4},$$

where the last inequality follows from $\|\mathbf{b}\|_2 = \|\tilde{\mathbf{b}}\|_2 = 1$. Appealing to Lemma B.1,

$$\liminf_{m \rightarrow \infty} \text{corr}(X, Y) \geq \frac{\|\mathbf{Q}_\perp^T \mathbf{b}\|_2^2 - \sqrt{6} \min\{\|\mathbf{Q}_\perp^T (\mathbf{b} \pm \tilde{\mathbf{b}})\|_2\}}{\|\mathbf{Q}_\perp^T \tilde{\mathbf{b}}\|_2^2}.$$

(3.11) follows by noting $\|\mathbf{Q}_\perp^T \mathbf{a}\|_2 = \|\mathbf{P}_{\mathbf{Q}_\perp} \mathbf{a}\|_2$ for $\mathbf{a} \in \mathbb{R}^N$.

To prove (3.13), we use Proposition 3.8 (i.e. (3.18)) to lower bound $\|\mathbf{Q}_\perp^T \tilde{\mathbf{b}}\|_2^2$:

$$\varphi \leq \|\mathbf{Q}_\perp^T \tilde{\mathbf{b}}\|_2 + \kappa \implies (\varphi - \kappa)^2 \leq \|\mathbf{Q}_\perp^T \tilde{\mathbf{b}}\|_2^2 \leq \|\tilde{\mathbf{b}}\|_2^2 = 1.$$

Also, $\|\mathbf{Q}_\perp^T \mathbf{b}\|_2^2 = 1 - \kappa^2$ and

$$\min\{\|\mathbf{Q}_\perp^T(\mathbf{b} \pm \tilde{\mathbf{b}})\|_2\} \leq \min\{\|\mathbf{b} \pm \tilde{\mathbf{b}}\|_2\} = \sqrt{2 - 2\varphi}.$$

Hence,

$$\liminf_{m \rightarrow \infty} \text{corr}(X, Y) \geq (1 - \kappa^2) - \frac{\sqrt{12(1 - \varphi)}}{(\varphi - \kappa)^2},$$

completing the proof. $\square$

C Proof of Theorem 3.11

We first prove the case of the sub-Gaussian sketches. According to Lemma 3.10, it suffices to verify the conditions (3.22).

Note $\sqrt{m}\mathbf{S}\mathbf{Q} \in \mathbb{R}^{m \times d}$ is a random matrix whose rows are i.i.d. isotropic random vectors in $\mathbb{R}^d$, with the sub-Gaussian norm $\lesssim K$ (this follows from Definition 3.4.1 and Proposition 2.6.1 in [Ver18]). Applying [Ver18, Theorem 4.6.1] to the matrix $\sqrt{m}\mathbf{S}\mathbf{Q}$ and using the fact that $\sigma_{\min}(\sqrt{m}\mathbf{S}\mathbf{Q}) = \sqrt{m}\sigma_{\min}(\mathbf{S}\mathbf{Q})$, we find that if $m \gtrsim K^4 d \log(4L/\delta)$, then with probability at least $1 - \delta/(2L)$, the first condition in (3.22) is satisfied.

For the second condition in (3.22), we write the i -th component of $\mathbf{Q}^T \mathbf{S}^T \mathbf{S} \mathbf{h}$ as

$$\mathbf{q}_i^T \mathbf{S}^T \mathbf{S} \mathbf{h} = \frac{1}{m} \sum_{j \in [m]} \langle \sqrt{m} \mathbf{s}_j, \mathbf{q}_i \rangle \langle \sqrt{m} \mathbf{s}_j, \mathbf{h} \rangle, \quad i \in [d]. \quad (\text{C.1})$$

Both $\langle \sqrt{m} \mathbf{s}_j, \mathbf{q}_i \rangle$ and $\langle \sqrt{m} \mathbf{s}_j, \mathbf{h} \rangle$ are sub-Gaussian random variables [Ver18, Proposition 2.6.1]. Therefore,

$$\|\langle \sqrt{m} \mathbf{s}_j, \mathbf{q}_i \rangle \langle \sqrt{m} \mathbf{s}_j, \mathbf{h} \rangle\|_{\psi_1} \leq \|\langle \sqrt{m} \mathbf{s}_j, \mathbf{q}_i \rangle\|_{\psi_2} \|\langle \sqrt{m} \mathbf{s}_j, \mathbf{h} \rangle\|_{\psi_2} \leq \|\sqrt{m} \mathbf{s}_j\|_{\psi_2}^2 \lesssim K^2,$$

where the first inequality follows from [Ver18, Lemma 2.7.7], the second inequality follows from [Ver18, Definition 3.4.1], and the final inequality follows from an application of [Ver18, Proposition 2.6.1]. Moreover, since $\mathbf{h} \perp \text{range}(\mathbf{Q})$ it is easy to verify that the summands in (C.1) are all zero-mean. By Bernstein's inequality [Ver18, Corollary 2.8.3], if $m \gtrsim K^4 d \log(4dL/\delta)/\varepsilon$, with probability at least $1 - \delta/(2dL)$, $|\mathbf{q}_i^T \mathbf{S}^T \mathbf{S} \mathbf{h}| \leq \sqrt{\varepsilon/(2d)}$. Taking a union bound over $i \in [d]$ yields that, with probability at least $1 - \delta/(2L)$,

$$\max_{i \in [d]} |\mathbf{q}_i^T \mathbf{S}^T \mathbf{S} \mathbf{h}| \leq \sqrt{\frac{\varepsilon}{2d}}. \quad (\text{C.2})$$

Note that (C.2) implies $\|\mathbf{Q}^T \mathbf{S}^T \mathbf{S} \mathbf{h}\|_2^2 \leq \frac{\varepsilon}{2}$. Consequently, combining the results we have that there exists an absolute constant C , such that if $m \geq CK^4 d \log(4dL/\delta)/\varepsilon$, then with probability at least $1 - \delta/L$,

$$\sigma_{\min}^2(\mathbf{S}\mathbf{Q}) \geq \frac{\sqrt{2}}{2} \quad \text{and} \quad \|\mathbf{Q}^T \mathbf{S}^T \mathbf{S} \mathbf{h}\|_2^2 \leq \frac{\varepsilon}{2},$$

which are the conditions in (3.22). This completes the proof for the sub-Gaussian sketch.

We next prove the case for the leverage score sampling matrices, and the proof is again based on Lemma 3.10. Note that leverage score sampling can be viewed as a special case of induced measure sampling. The first condition in (3.22) is implied by $\|\mathbf{Q}^T \mathbf{S}^T \mathbf{S} \mathbf{Q} - \mathbf{I}\|_2 \leq 1 - \frac{\sqrt{2}}{2}$, which, according to [Mal+22, Lemma A.1], is satisfied with probability at least $1 - \delta/2L$ if $m \geq 35d \log(4dL/\delta)$. For the second condition in (3.22), the only difference is that one uses Markov's inequality in place of Bernstein's inequality due to the lack of information on the tail of $\mathbf{q}_i^T \mathbf{S}^T \mathbf{S} \mathbf{h}$, and the details are omitted. Under additional assumptions in (3.25), Markov's inequality can be replaced by Hoeffding's inequality to yield an improved bound (3.26):

$$\mathbf{q}_i^T \mathbf{S}^T \mathbf{S} \mathbf{h} = \frac{1}{m} \sum_{j \in [m]} \langle \sqrt{m} \mathbf{s}_j, \mathbf{q}_i \rangle \langle \sqrt{m} \mathbf{s}_j, \mathbf{h} \rangle, \quad i \in [d],$$

with each summand $\langle \sqrt{m} \mathbf{s}_j, \mathbf{q}_i \rangle \langle \sqrt{m} \mathbf{s}_j, \mathbf{h} \rangle$ centered and bounded as

$$|\langle \sqrt{m} \mathbf{s}_j, \mathbf{q}_i \rangle \langle \sqrt{m} \mathbf{s}_j, \mathbf{h} \rangle| \leq \max_{i \in [d]} \max_{j \in [N]: \ell_j > 0} \frac{r |q_{ij} h_j|}{\ell_j} \leq \max_{i \in [d]} \max_{j \in [N]: \ell_j > 0} \frac{d |q_{ij} h_j|}{\ell_j} \leq C,$$

where r is the rank of $\mathbf{A}$. By Hoeffding's inequality, for $t > 0$,

$$\mathbb{P}(|\mathbf{q}_i^T \mathbf{S}^T \mathbf{S} \mathbf{h}| \leq t) \geq 1 - 2 \exp\left(-\frac{mt^2}{2C^2}\right).$$

Setting $t = \sqrt{\varepsilon/2d}$ and taking a union bound over i yields that, for $m \geq 4C^2 d \log(4dL/\delta)/\varepsilon$, with probability at least $1 - \delta/2L$, $\|\mathbf{Q}^T \mathbf{S}^T \mathbf{S} \mathbf{h}\|_2^2 \leq \varepsilon/2$.

References

- [ABW22] B. Adcock, S. Brugiapaglia, and C. G. Webster. *Sparse Polynomial Approximation of High-Dimensional Functions*. Vol. 25. SIAM, 2022.
- [BC15] M. Bachmayr and A. Cohen. *Kolmogorov widths and low-rank approximations of parametric elliptic PDEs*. 2015. DOI: [10.48550/ARXIV.1502.03117](https://doi.org/10.48550/ARXIV.1502.03117).
- [De+20] S. De, J. Britton, M. Reynolds, R. Skinner, K. Jansen, and A. Doostan. “On transfer learning of neural networks using bi-fidelity data for uncertainty propagation”. *International Journal for Uncertainty Quantification* 10.6 (2020).
- [DD22] S. De and A. Doostan. “Neural network training using ℓ_1 -regularization and bi-fidelity data”. *Journal of Computational Physics* 458 (2022), p. 111010.
- [DW17] M. Dereziński and M. K. Warmuth. “Unbiased Estimates for Linear Regression via Volume Sampling”. *arXiv preprint arXiv:1705.06908* (2017). arXiv: [1705.06908](https://arxiv.org/abs/1705.06908).
- [DW18] M. Dereziński and M. K. Warmuth. “Reverse Iterative Volume Sampling for Linear Regression”. *The Journal of Machine Learning Research* 19.1 (2018), pp. 853–891.
- [DWH18] M. Dereziński, M. K. Warmuth, and D. Hsu. “Leveraged Volume Sampling for Linear Regression”. *arXiv preprint arXiv:1802.06749* (2018). arXiv: [1802.06749](https://arxiv.org/abs/1802.06749).
- [DDH18] P. Diaz, A. Doostan, and J. Hampton. “Sparse polynomial chaos expansions via compressed sensing and D-optimal design”. *Computer Methods in Applied Mechanics and Engineering* 336 (2018), pp. 640–666.
- [DGRH07] A. Doostan, R. G. Ghanem, and J. Red-Horse. “Stochastic model reduction for chaos representations”. *Computer Methods in Applied Mechanics and Engineering* 196.37 (2007), pp. 3951–3966. ISSN: 0045-7825. DOI: <https://doi.org/10.1016/j.cma.2006.10.047>.
- [DO11] A. Doostan and H. Owhadi. “A non-adapted sparse approximation of PDEs with stochastic inputs”. *Journal of Computational Physics* 230.8 (2011), pp. 3015–3034.
- [Dri+12] P. Drineas, M. Magdon-Ismail, M. W. Mahoney, and D. P. Woodruff. “Fast Approximation of Matrix Coherence and Statistical Leverage”. *The Journal of Machine Learning Research* 13.1 (2012), pp. 3475–3506.
- [DMM08] P. Drineas, M. W. Mahoney, and S. Muthukrishnan. “Relative-Error CUR Matrix Decompositions”. *SIAM Journal on Matrix Analysis and Applications* 30.2 (2008), pp. 844–881.
- [DMM06] P. Drineas, M. W. Mahoney, and S. Muthukrishnan. “Sampling Algorithms for ℓ_2 Regression and Applications”. In: *Proceedings of the Seventeenth Annual ACM-SIAM Symposium on Discrete Algorithms*. 2006, pp. 1127–1136.
- [Dri+11] P. Drineas, M. W. Mahoney, S. Muthukrishnan, and T. Sarlós. “Faster Least Squares Approximation”. *Numerische Mathematik* 117.2 (2011), pp. 219–249.
- [Fai+17] H. R. Fairbanks, A. Doostan, C. Ketelsen, and G. Iaccarino. “A low-rank control variate for multilevel Monte Carlo simulation of high-dimensional uncertain systems”. *Journal of Computational Physics* 341 (2017), pp. 121–139. ISSN: 0021-9991. DOI: <https://doi.org/10.1016/j.jcp.2017.03.060>.

- [Fai+20] H. R. Fairbanks, L. Jofre, G. Geraci, G. Iaccarino, and A. Doostan. “Bi-fidelity approximation for uncertainty quantification and sensitivity analysis of irradiated particle-laden turbulence”. *Journal of Computational Physics* 402 (2020), p. 108996.
- [GS03] R. G. Ghanem and P. D. Spanos. *Stochastic finite elements: a spectral approach*. Courier Corporation, 2003.
- [GVL13] G. H. Golub and C. F. Van Loan. *Matrix Computations*. English. Fourth. Baltimore: Johns Hopkins University Press, 2013. ISBN: 978-1-4214-0794-4.
- [Guo+18] L. Guo, A. Narayan, L. Yan, and T. Zhou. “Weighted Approximate Fekete Points: Sampling for Least-Squares Polynomial Approximation”. *SIAM Journal on Scientific Computing* 40.1 (2018). arXiv:1708.01296 [math.NA], A366–A387. ISSN: 1064-8275. DOI: [10.1137/17M1140960](https://doi.org/10.1137/17M1140960).
- [HNP22] C. Haberstick, A. Nouy, and G. Perrin. “Boosted optimal weighted least-squares”. *Mathematics of Computation* 91.335 (2022), pp. 1281–1315. DOI: [10.1090/mcom/3710](https://doi.org/10.1090/mcom/3710).
- [HD18a] M. Hadigol and A. Doostan. “Least squares polynomial chaos expansion: A review of sampling strategies”. *Computer Methods in Applied Mechanics and Engineering* 332 (2018), pp. 382–407.
- [HMT11] N. Halko, P.-G. Martinsson, and J. A. Tropp. “Finding Structure with Randomness: Probabilistic Algorithms for Constructing Approximate Matrix Decompositions”. *SIAM Review* 53.2 (May 2011), pp. 217–288.
- [HD15a] J. Hampton and A. Doostan. “Coherence motivated sampling and convergence analysis of least squares polynomial Chaos regression”. *Computer Methods in Applied Mechanics and Engineering* 290 (2015), pp. 73–97. ISSN: 0045-7825. DOI: <https://doi.org/10.1016/j.cma.2015.02.006>.
- [HD15b] J. Hampton and A. Doostan. “Compressive sampling of polynomial chaos expansions: Convergence analysis and sampling strategies”. *Journal of Computational Physics* 280 (2015), pp. 363–386. ISSN: 0021-9991. DOI: <https://doi.org/10.1016/j.jcp.2014.09.019>.
- [HD18b] J. Hampton and A. Doostan. “Basis adaptive sample efficient polynomial chaos (BASE-PC)”. *Journal of Computational Physics* 371 (2018), pp. 20–49. ISSN: 0021-9991. DOI: <https://doi.org/10.1016/j.jcp.2018.03.035>.
- [Ham+18a] 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. ISSN: 0021-9991. DOI: <https://doi.org/10.1016/j.jcp.2018.04.015>.
- [Ham+18b] 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.
- [Kol05] T. Kollo. *Advanced multivariate statistics with matrices*. Springer, 2005.
- [LK20] B. W. Larsen and T. G. Kolda. “Practical Leverage-Based Sampling for Low-Rank Tensor Decomposition”. *arXiv preprint arXiv:2006.16438v3* (2020). arXiv: [arXiv:2006.16438v3](https://arxiv.org/abs/2006.16438v3).
- [LMK10] O. Le Maître and O. M. Knio. *Spectral methods for uncertainty quantification: with applications to computational fluid dynamics*. Springer Science & Business Media, 2010.
- [Mah11] M. W. Mahoney. “Randomized Algorithms for Matrices and Data”. *Foundations and Trends in Machine Learning* 3.2 (2011), pp. 123–224.
- [Mal+22] O. A. Malik, Y. Xu, N. Cheng, S. Becker, A. Doostan, and A. Narayan. *Fast algorithms for monotone lower subsets of Kronecker least squares problems*. Preprint. 2022.
- [MT20] P.-G. Martinsson and J. A. Tropp. “Randomized numerical linear algebra: Foundations and algorithms”. *Acta Numerica* 29 (2020), 403–572. DOI: [10.1017/S0962492920000021](https://doi.org/10.1017/S0962492920000021).
- [NGX14] A. Narayan, C. Gittelsohn, and D. Xiu. “A Stochastic Collocation Algorithm with Multifidelity Models”. *SIAM Journal on Scientific Computing* 36.2 (2014), A495–A521. ISSN: 1064-8275, 1095-7197. DOI: [10.1137/130929461](https://doi.org/10.1137/130929461).
- [New+22] F. Newberry, J. Hampton, K. Jansen, and A. Doostan. “Bi-fidelity reduced polynomial chaos expansion for uncertainty quantification”. *Computational Mechanics* 69.2 (2022), pp. 405–424.

- [PWG18] B. Peherstorfer, K. Willcox, and M. Gunzburger. “Survey of Multifidelity Methods in Uncertainty Propagation, Inference, and Optimization”. *SIAM Review* 60.3 (2018), pp. 550–591. ISSN: 0036-1445. DOI: [10.1137/16M1082469](https://doi.org/10.1137/16M1082469).
- [PHD14] J. Peng, J. Hampton, and A. Doostan. “A weighted ℓ_1 -minimization approach for sparse polynomial chaos expansions”. *Journal of Computational Physics* 267 (2014), pp. 92–111. ISSN: 0021-9991. DOI: <https://doi.org/10.1016/j.jcp.2014.02.024>.
- [SNM17] P. Seshadri, A. Narayan, and S. Mahadevan. “Effectively Subsampled Quadratures for Least Squares Polynomial Approximations”. *SIAM/ASA Journal on Uncertainty Quantification* 5.1 (2017), pp. 1003–1023.
- [Smi13] R. C. Smith. *Uncertainty quantification: theory, implementation, and applications*. Vol. 12. Siam, 2013.
- [Ver18] R. Vershynin. *High-Dimensional Probability: An Introduction with Applications in Data Science*. Vol. 47. Cambridge University Press, 2018.
- [Wic50] G.-C. Wick. “The evaluation of the collision matrix”. *Physical review* 80.2 (1950), p. 268.
- [Woo14] D. P. Woodruff. “Sketching as a Tool for Numerical Linear Algebra”. *Foundations and Trends in Theoretical Computer Science* 10.1-2 (2014), pp. 1–157.
- [XH05] D. Xiu and J. S. Hesthaven. “High-Order Collocation Methods for Differential Equations with Random Inputs”. *SIAM Journal on Scientific Computing* 27.3 (2005), pp. 1118–1139. DOI: [10.1137/040615201](https://doi.org/10.1137/040615201).
- [XK02] D. Xiu and G. E. Karniadakis. “The Wiener–Askey Polynomial Chaos for Stochastic Differential Equations”. *SIAM Journal on Scientific Computing* 24.2 (2002), pp. 619–644. DOI: [10.1137/S1064827501387826](https://doi.org/10.1137/S1064827501387826).
- [ZNX14] X. Zhu, A. Narayan, and D. Xiu. “Computational Aspects of Stochastic Collocation with Multifidelity Models”. *SIAM/ASA Journal on Uncertainty Quantification* 2.1 (2014), pp. 444–463. DOI: [10.1137/130949154](https://doi.org/10.1137/130949154).