

SUPPLEMENTARY MATERIALS: Confidence intervals for functionals in constrained inverse problems via data-adaptive sampling-based calibration

Michael Stanley*, Pau Batlle , Pratik Patil†, Houman Owhadi‡, and Mikael Kuusela§

This document serves as a supplement to the paper “Confidence intervals for functionals in constrained inverse problems via data-adaptive sampling-based calibration”. Below we provide an organization for the supplement, followed by a summary of the main notation used in both the paper and the appendix.

Organization. A roadmap for the supplement is provided in [Table SM1](#).

Table SM1: Roadmap of the supplement.

Appendix	Description
Section SM1	Proofs in Section 3 (Proofs of Lemma 3.1 and Corollary 3.2)
Section SM2	Proofs in Section 4 (Proof of Theorem 4.1)
Section SM3	Additional details and illustrations in Subsection 5.1 (VGS and Polytope samplers)
Section SM4	Additional details in Subsection 5.2 (quantile regression)
Section SM5	Additional details and illustrations in Section 6 (importance-like sampler)

Notation. An overview of the main notation used in this paper is provided in [Table SM2](#).

SM1. Proofs in [Section 3](#).

SM1.1. Proof of [Lemma 3.1](#). We reproduce the original argument in [\[SM2\]](#), originally stated in terms of p-values, translated into our quantile setting. Fix any $\mathbf{x}^* \in \mathcal{X}$ and consider the sets:

- $A_1 = \{\mathbf{y} : B_\eta(\mathbf{y}) \ni \mathbf{x}^*\}$
- $A_2 = \{\mathbf{y} : \lambda(\mu^*, \mathbf{y}) \leq Q_{\mathbf{x}}(1 - \gamma)\}$
- $A_3 = \{\mathbf{y} : \lambda(\mu^*, \mathbf{y}) \leq \bar{q}_{\gamma, \eta}(\mu^*)\} = \{\mathbf{y} : \lambda(\mu^*, \mathbf{y}) \leq \sup_{\mathbf{x} \in \mathcal{B}_\eta \cap \Phi_{\mu^*}} Q_{\mathbf{x}}(1 - \gamma)\}$

We now have that $\mathbb{P}(\mathbf{y} \in A_1) = 1 - \eta$, $\mathbb{P}(\mathbf{y} \in A_2) = 1 - \gamma$, and that $(A_1 \cap A_2) \subset (A_1 \cap A_3)$. Therefore,

$$(SM1.1) \quad \mathbb{P}(\mathbf{y} \notin A_3) = \mathbb{P}(\mathbf{y} \notin A_3, \mathbf{y} \in A_1) + \mathbb{P}(\mathbf{y} \notin A_3, \mathbf{y} \notin A_1)$$

$$(SM1.2) \quad \leq \mathbb{P}(\mathbf{y} \notin A_2, \mathbf{y} \in A_1) + \mathbb{P}(\mathbf{y} \notin A_1)$$

*Analytical Mechanics Associates, USA

†Department of Statistics and Data Sciences, University of Texas at Austin, USA

‡Department of Computing and Mathematical Sciences, California Institute of Technology, USA

§Department of Statistics and Data Science, Carnegie Mellon University, USA

Table SM2: Summary of main notation used in the paper and the appendix.

Notation	Description
Non-bold lower or upper case	Denotes scalars (e.g., α , μ , Q).
Bold lower case	Denotes vectors (e.g., $\mathbf{x}$, $\mathbf{y}$, $\mathbf{h}$).
Bold upper case	Denotes matrices (e.g., $\mathbf{K}$, $\mathbf{I}$).
Calligraphic font	Denotes sets (e.g., $\mathcal{X}$, $\mathcal{C}$, $\mathcal{D}$).
$\mathbb{R}$	Set of real numbers.
$\mathbb{R}_+$	Set of non-negative real numbers.
$[n]$	Set $\{1, \dots, n\}$ for a positive integer n .
$\ \mathbf{u}\ _2$	The ℓ_2 norm of vector $\mathbf{u}$.
$\ f\ _{L_2}$	The L_2 norm of function f .
$\mathbf{K}^\top$	The transpose of a matrix $\mathbf{K} \in \mathbb{R}^{m \times p}$.
$\mathbf{I}_m$ or $\mathbf{I}$	The $m \times m$ identity matrix.
$\mathbf{v} \leq \mathbf{u}$	Componentwise ordering for vectors $\mathbf{v}$ and $\mathbf{u}$.
$\mathbf{A} \preceq \mathbf{B}$	The Loewner ordering for symmetric matrices $\mathbf{A}$ and $\mathbf{B}$.
$\mathbf{1}\{A\}$	Indicator random variable associated with event A .
$Y = \mathcal{O}_\alpha(X)$	Deterministic big-O notation, indicating that Y is bounded by $ Y \leq C_\alpha X$.
C_α	A numerical constant that may depend on the parameter α in context.
$\mathcal{O}_p$	Probabilistic big-O notation.
$\xrightarrow{\mathbb{P}}$	Convergence in probability.
$X \succeq Y$	Stochastic dominance of X by Y , indicating that $\mathbb{P}(X \geq z) \geq \mathbb{P}(Y \geq z)$ for all $z \in \mathbb{R}$.

$$(SM1.3) \quad \leq \mathbb{P}(\mathbf{y} \notin A_2) + \mathbb{P}(\mathbf{y} \notin A_1)$$

$$(SM1.4) \quad = \gamma + \eta$$

so that $\mathbb{P}(\mathbf{y} \in A_3) \geq 1 - \gamma - \eta$. Imposing $1 - \gamma - \eta \geq 1 - \alpha$ gives the desired result.

SM1.2. Proof of Corollary 3.2. By definition, $\bar{q}_{\gamma,\eta}^\mu \leq \bar{q}_{\gamma,\eta}$ for all $\mu \in \mathbb{R}$. Therefore, $C_\alpha^{\text{sl}}(\mathbf{y}; \mathcal{B}_\eta) \subseteq C_\alpha^{\text{gl}}(\mathbf{y}; \mathcal{B}_\eta)$, and thus

$$(SM1.5) \quad \mathbb{P}_{\mathbf{x}^*}(\mu^* \in C_\alpha^{\text{gl}}(\mathbf{y}; \mathcal{B}_\eta)) \geq \mathbb{P}_{\mathbf{x}^*}(\mu^* \in C_\alpha^{\text{sl}}(\mathbf{y}; \mathcal{B}_\eta)) \geq 1 - \alpha.$$

SM2. Proof of Theorem 4.1. Throughout the proof, we make use of the following lemma:

Lemma SM2.1. Let $f : \mathbb{R}^n \rightarrow \mathbb{R}$ and $\mathcal{X} \subset \mathbb{R}^n$ such that $\min_{\mathbf{x} \in \mathcal{X}} f(\mathbf{x})$ is achieved, and at least one of the minimizers $\mathbf{x}^*$ satisfies:

1. f is continuous at $\mathbf{x}^*$,
2. $\mathbf{x}^*$ is not an isolated point of $\mathcal{X}$ (i.e., $\forall \delta > 0, B_\delta(\mathbf{x}^*) \cap \mathcal{X} \neq \emptyset$).

Let μ be a measure on $\mathcal{X}$ such that $\mu(B) > 0$ for all $B \subseteq \mathcal{X}$ such that $\lambda_{\text{Leb}}(B) > 0$ (where $\lambda_{\text{Leb}}(B)$ refers here to the Lebesgue measure of the set B). Let $Y_m = \min_{i=1, \dots, m} f(\mathbf{x}_i)$, where $\mathbf{x}_i$

are i.i.d. samples from μ . Then $Y_m \xrightarrow{P} f(\mathbf{x}^*)$.

Proof. Fix $\varepsilon > 0$ and let us show that $\mathbb{P}(|Y_m - f(\mathbf{x}^*)| > \varepsilon) \rightarrow 0$. Since f is continuous at $\mathbf{x}^*$ there exists a $\delta > 0$ such that $f(B_\delta(\mathbf{x}^*)) \subset B_\varepsilon(f(\mathbf{x}^*))$ so that,

$$(SM2.1) \quad \mathbb{P}(|Y_m - f(\mathbf{x}^*)| \geq \varepsilon) \leq \mathbb{P}(\mathbf{x}_i \notin B_\delta(\mathbf{x}^*), \forall i = 1, \dots, m) = (\mathbb{P}_{\mathbf{x} \sim \mu}(\mathbf{x} \notin B_\delta(\mathbf{x}^*)))^m.$$

Since $\mathbf{x}^*$ is not an isolated point, we have $\lambda_{\text{Leb}}(B_\delta(\mathbf{x}^*) \cap \mathcal{X}) > 0$ and therefore $\mathbb{P}_{\mathbf{x} \sim \mu}(\mathbf{x} \notin B_\delta(\mathbf{x}^*)) < 1$ and $\mathbb{P}(|Y_m - f(\mathbf{x}^*)| \geq \varepsilon) \rightarrow 0$. ■

We begin by proving that the empirical maximums of the quantiles obtained both by [Algorithms 3.1](#) and [3.2](#) (assuming the quantile regressor is consistent) converge to the true max quantile. **Algorithm 3.1**

Let $\bar{q}_{\gamma, \eta}^{\text{de}} := \max_{i=1, \dots, M} \hat{q}_\gamma^i(N)$, where we explicitly write the dependence with the number of samples and the index i refers to the quantile estimated at the i -th sampled point $\mathbf{x}_i$. We aim to show that $\bar{q}_{\gamma, \eta}^{\text{de}} \xrightarrow{P} \bar{q}_{\gamma, \eta}$. We know that $\hat{q}_\gamma^i(N)$ converges in probability to $Q_{P_{\mathbf{x}_i}}(1 - \gamma)$ as $N \rightarrow \infty$. We have

$$(SM2.2) \quad \left| \bar{q}_{\gamma, \eta}^{\text{de}} - \bar{q}_{\gamma, \eta} \right| \leq \left| \bar{q}_{\gamma, \eta}^{\text{de}} - \max_{i=1, \dots, M} Q_{P_{\mathbf{x}_i}}(1 - \gamma) \right| + \left| \max_{i=1, \dots, M} Q_{P_{\mathbf{x}_i}}(1 - \gamma) - \bar{q}_{\gamma, \eta} \right|.$$

The first term can be made smaller than $\varepsilon/2$ as $N \rightarrow \infty$ by convergence of the estimator, and the second term can be made smaller than $\varepsilon/2$ as $M \rightarrow \infty$ by application of [Lemma SM2.1](#) to the quantile function (maximizing instead of minimizing).

Algorithm 3.2

The proof is identical to that of [Algorithm 3.1](#), with the only difference of replacing the quantiles estimated via Monte Carlo sampling with those estimated by the quantile regression, and N to M_{tr} , the number of samples needed to train the quantile regression. Since the quantile regression is assumed to be consistent, the first term can be made arbitrarily small as M_{tr} grows, and the result follows.

Since identical convergence results apply for both algorithms, we will not explicitly distinguish. The proof is written in terms of N , which can be replaced by M_{tr} .

SM2.1. Proof of Statement 1 (Global Inverted). Recall,

$$(SM2.3) \quad C_{\text{inv}}^{\text{gl}}(\mathbf{y}) = \left[\min_{k \in \{1, \dots, M\} : \lambda(\varphi(\mathbf{x}_k), \mathbf{y}) \leq \hat{q}(N)} \varphi(\mathbf{x}_k), \max_{k : \lambda(\varphi(\mathbf{x}_k), \mathbf{y}) \leq \hat{q}(N)} \varphi(\mathbf{x}_k) \right]$$

$$(SM2.4) \quad = \left[\min_{k \in \{1, \dots, M\} : \lambda(\mu_k, \mathbf{y}) \leq \hat{q}(N)} \mu_k, \max_{k : \lambda(\mu_k, \mathbf{y}) \leq \hat{q}(N)} \mu_k \right]$$

where $\hat{q}(N) := \max_{i=1, \dots, M} \hat{q}_\gamma^i(N)$, the estimated quantiles of the sampled $\mathbf{x}_i \in \mathcal{B}_\eta$ and $\mu_i := \varphi(\mathbf{x}_i)$.

Note that μ_i are samples in $\varphi(\mathcal{B}_\eta) \subset \mathbb{R}$. Also note that we have more explicitly written out the interval definition (i.e., [Equation \(3.8\)](#)) to emphasize clarity rather than presentation.

Consider the left extreme of the interval; a similar argument follows from the right extreme. Consider three quantities:

$$(SM2.5) \quad \mu_1(N, M) := \min \mu_k \quad \text{s.t.} \quad k = 1, \dots, M \quad \text{and} \quad \lambda(\mu_k, \mathbf{y}) \leq \hat{q}(N)$$

$$(SM2.6) \quad \mu_2(M) := \min \mu_k \quad \text{s.t.} \quad k = 1, \dots, M \text{ and } \lambda(\mu_k, \mathbf{y}) \leq \bar{q}_{\gamma, \eta}$$

$$(SM2.7) \quad \mu_3 := \min \mu \quad \text{s.t.} \quad \mu \in \varphi(\mathcal{B}_\eta) \text{ and } \lambda(\mu, \mathbf{y}) \leq \bar{q}_{\gamma, \eta}$$

Our goal is to show that as $N, M \rightarrow \infty$, $\mu_1 \xrightarrow{P} \mu_3$. Lemma SM2.1 shows that $\mu_2 \xrightarrow{P} \mu_3$. Indeed, the minimization over the indices k such that the condition is satisfied can be seen as a rejection sampling strategy in which all accepted samples are samples of the feasible region of the optimization in μ_3 . As M grows, since the sampler eventually samples all areas of $\varphi(\mathcal{B}_\eta)$, some samples are guaranteed to be close to the optimum with high probability. Finally, for fixed M and N going to infinity, $\mu_1 \xrightarrow{P} \mu_2$. This follows from the continuity of the optimization problem with respect to the right-hand side of the constraint, and the fact that $\max_{i=1, \dots, M} q(\mathbf{x}_i) \xrightarrow{P} \bar{q}_{\gamma, \eta}$. It follows that as $N, M \rightarrow \infty$, $\mu_1 \xrightarrow{P} \mu_3$.

SM2.2. Proof of Statement 2 (Sliced Inverted). The proof technique is similar to the one of Statement 1, replacing $\hat{q}(N) := \max_{i=1, \dots, M} \hat{q}_\gamma^i(N)$ for $\hat{q}_\gamma^k(N)$. Defined then, similarly to the previous proof:

$$(SM2.8) \quad \mu_1(N, M) := \min \mu_k \quad \text{s.t.} \quad k = 1, \dots, M \text{ and } \lambda(\mu_k, \mathbf{y}) \leq \hat{q}_\gamma^k(N)$$

$$(SM2.9) \quad \mu_2(M) := \min \mu_k \quad \text{s.t.} \quad k = 1, \dots, M \text{ and } \lambda(\mu_k, \mathbf{y}) \leq \bar{q}_{\gamma, \eta}(\mu_k)$$

$$(SM2.10) \quad \mu_3 := \min \mu \quad \text{s.t.} \quad \mu \in \varphi(\mathcal{B}_\eta) \text{ and } \lambda(\mu, \mathbf{y}) \leq \bar{q}_{\gamma, \eta}(\mu)$$

where $\bar{q}_{\gamma, \eta}(\mu) = \max_{\mathbf{x} \in \Phi_\mu \cap \mathcal{B}_\eta} Q_{\mathbf{x}}(1 - \gamma)$. As $N \rightarrow \infty$, $\hat{q}_\gamma^k(N) \rightarrow Q_{\mathbf{x}_k}(1 - \gamma)$. Lemma SM2.1 can be used to show $\mu_2 \xrightarrow{P} \mu_3$, because the sampler strategy used that first samples $\mathbf{x}_k$ and then accepts $\mu_k = \varphi(\mathbf{x}_k)$ as a sample if $\lambda(\varphi(\mathbf{x}_k), \mathbf{y}) < q_\gamma^k$, it can be shown that for every feasible point μ , there is eventually a sample close to it.

Finally, μ_1 will become arbitrarily close to μ_2 as M grows large, since both are taking the minimum over samples that are sampled densely from the feasible set, meaning accepted μ_k will eventually be close to μ_3 for both the case of μ_1 and the case of μ_2 .

SM2.3. Proof of Statement 3 (Global Optimized). We prove the continuity of the optimization problem,

$$(SM2.11) \quad \max_{\mu \in \varphi(\mathcal{B}_\eta)} \mu \quad \text{s.t.} \quad \lambda(\mu, \mathbf{y}) \leq q,$$

as a function of q in the positive measure interval $(\lambda(\bar{\mu}, \mathbf{y}), \bar{q}_{\gamma, \eta}]$. Therefore, convergence in probability follows as we have convergence in probability to $\bar{q}_{\gamma, \eta}$ as $N, M \rightarrow \infty$, which in particular implies that the maximum quantile estimate is in the interval $(\lambda(\bar{\mu}, \mathbf{y}), \bar{q}_{\gamma, \eta}]$ almost surely. We do so by appealing to the maximum theorem [SM16, § E.3], which, in general, guarantees continuity of functions of the form $f^*(\theta) = \sup\{f(x, \theta) : x \in C(\theta)\}$ as long as f is continuous, C is a continuous compact-valued correspondence and $C(\theta)$ is non-empty for all $\theta \in \Theta$. The continuity of C comes from the continuity of the LLR, f is equal to the identity and the strict feasibility condition ensures C is non-empty in $\Theta := (\lambda(\bar{\mu}, \mathbf{y}), \bar{q}_{\gamma, \eta}]$.

SM2.4. Proof of Statement 3 (Sliced Optimized). We will prove sufficient conditions for convergence of $\inf_{\mu: \hat{f}_k(\mu) \geq 0} \mu$ to $\inf_{\mu: f(\mu) \geq 0} \mu$ as $\hat{f}_k$ converges to f , and the result will follow by

taking $\widehat{f}_k(\mu) = \widehat{m}_\gamma(\mu) - \lambda(\mu, \mathbf{y})$ and $f(\mu) = m_\gamma(\mu) - \lambda(\mu, \mathbf{y})$. A similar argument can be repeated for the supremum. Use the notation $\widehat{f}_k$ to indicate that k sampled points are used to estimate this function via the definition of $\widehat{m}_\gamma(\mu)$ (see [Subsection 3.2](#)).

Lemma SM2.2. *Let $f : \mathbb{R} \rightarrow \mathbb{R}$ be a function, and let f_k be a sequence of functions $f_k : \mathbb{R} \rightarrow \mathbb{R}$. Let $\mu^* = \inf_{f(\mu) \geq 0} \mu$ and $\mu_k = \inf_{f_k(\mu) \geq 0} \mu$. Let the sequence of functions $\{f_k\}$ be such that for all $\delta > 0$,*

$$(SM2.12) \quad \mathbb{P} \left(\sup_{\mu} |f_k(\mu) - f(\mu)| > \delta \right) \rightarrow 0 \text{ as } k \rightarrow \infty,$$

namely, f_k converges in probability uniformly to f . Furthermore, let f be such that for all $\varepsilon > 0$, there exists $\delta' > 0$ such that if $|\mu - \mu^| \geq \varepsilon$, then $|f(\mu)| > \delta'$. Then, we have $\mu_k \xrightarrow{P} \mu^*$.*

Proof. Assume for the sake of contradiction that there exists $\varepsilon > 0$ such that $\mathbb{P}(|\mu_k - \mu^*| \geq \varepsilon)$ does not go to 0. Then, by the condition on f , it must follow that $\mathbb{P}(|f(\mu_k)| > \delta)$ does not go to 0. But,

$$(SM2.13) \quad \mathbb{P}(|f(\mu_k)| > \delta) \leq \mathbb{P}(|f(\mu_k) - f_k(\mu_k)| > \delta) + \mathbb{P}(|f_k(\mu_k)| > \delta),$$

and the right-hand side goes to 0 by uniform convergence in probability (first term) and by feasibility of μ_k (second term). ■

SM3. Sampling the Berger-Boos set.

SM3.1. VGS Sampler for low-dimensional and full column rank settings. Under the linear-Gaussian assumptions, when $\mathbf{K}$ has full column rank, it can be shown that

$$(SM3.1) \quad \mathcal{B}_\eta = \{\mathbf{x} \in \mathcal{X} : (\mathbf{x} - \widehat{\mathbf{x}})^\top \mathbf{K}^\top \boldsymbol{\Sigma}^{-1} \mathbf{K} (\mathbf{x} - \widehat{\mathbf{x}}) \leq \chi_{n,\eta}^2\},$$

where $\widehat{\mathbf{x}}$ is the Generalized Least-Squares (GLS) estimator. Equation (SM3.1) describes an ellipsoid in $\mathbb{R}^p$ intersected with $\mathcal{X}$ with axis directions and lengths determined by $\chi_{n,\eta}^2$ and the eigenvectors and eigenvalues of $\mathbf{K}^\top \boldsymbol{\Sigma}^{-1} \mathbf{K}$, respectively (see [SM17] for more details). An ellipsoid is nothing but a deformed ball. As such, we can sample uniformly at random from the ellipsoid in $\mathbb{R}^p$ using an algorithm sampling uniformly at random from a p -ball, followed by an appropriate linear transformation and translation. We can then include an additional accept-reject step to account for the constraint set $\mathcal{X}$. Note that all subsequent discussions of spheres and balls assume unit radii and centering at the origin. Also note that because the ellipsoid in (SM3.1) is centered around the GLS estimator, we can sample points first from within an ellipsoid of the same shape centered at the origin, and then translate those points by $\widehat{\mathbf{x}}$.

[SM23] propose a particularly efficient and clever algorithm to sample uniformly at random from the p -ball by proving a connection between uniform sampling on the $(p+1)$ -sphere (i.e., the surface of the ball in $\mathbb{R}^{p+2}$) and uniform sampling within the p -ball (i.e., within the unit ball in $\mathbb{R}^p$). Namely, one can sample a Gaussian $\tilde{\mathbf{z}} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_{p+2})$ and normalize $\mathbf{z} := \tilde{\mathbf{z}} / \|\tilde{\mathbf{z}}\|_2$ to sample a point uniformly at random from the $(p+1)$ -sphere. Next, one simply drops the last two elements of $\mathbf{z}$ to obtain a point that is uniformly sampled from within a p -ball. The

validity of this approach is substantiated by Lemma 1 and Theorem 1 in [SM23] and relies upon a distribution result about the ratio of chi-squared distributions and preservation of distribution under an orthogonal transformation. To denote sampling a point following this procedure, we use the notation $\mathbf{x} \sim \text{VGS}(p)$.

To sample from our desired ellipsoid in (SM3.1), first consider the eigendecomposition of $\mathbf{K}^\top \boldsymbol{\Sigma}^{-1} \mathbf{K} = \mathbf{P} \boldsymbol{\Omega} \mathbf{P}^\top$, where $\boldsymbol{\Omega} = \text{diag}(\omega_1^2, \omega_2^2, \dots, \omega_p^2)$, ω_i is the i -th eigenvalue of $\mathbf{K}^\top \boldsymbol{\Sigma}^{-1} \mathbf{K}$ and $\mathbf{P}$ is an orthonormal matrix where the columns vectors are the eigenvectors of $\mathbf{K}^\top \boldsymbol{\Sigma}^{-1} \mathbf{K}$. Denote $\boldsymbol{\Omega}^{1/2} = \text{diag}(\omega_1, \omega_2, \dots, \omega_p)$ and by extension, $\boldsymbol{\Omega}^{-1/2} = \text{diag}(\omega_1^{-1}, \omega_2^{-1}, \dots, \omega_p^{-1})$ when $\omega_i > 0$ for all i . These decompositions imply the correct transformation to apply to points sampled from the p -ball via $\mathbf{x} \sim \text{VGS}(p)$. Namely, define $\mathbf{w} := \sqrt{\chi_{n,\eta}^2} \mathbf{P} \boldsymbol{\Omega}^{-1/2} \mathbf{x}$. To know that $\mathbf{w}$ is sampled from the correct ellipsoid, it should be the case that $\mathbf{w}^\top \mathbf{K}^\top \boldsymbol{\Sigma}^{-1} \mathbf{K} \mathbf{w} \leq \chi_{n,\eta}^2$, as then we can simply make the update $\mathbf{w} \leftarrow \mathbf{w} + \hat{\mathbf{x}}$ to ensure that we have a point sampled from (SM3.1). This guarantee is verified as follows:

$$\begin{aligned}
 \mathbf{w}^\top \mathbf{K}^\top \boldsymbol{\Sigma}^{-1} \mathbf{K} \mathbf{w} &= \mathbf{w}^\top \mathbf{P} \boldsymbol{\Omega}^{1/2} \boldsymbol{\Omega}^{1/2} \mathbf{P}^\top \mathbf{w} \\
 (\text{SM3.2}) \quad &= \chi_{n,\eta}^2 \cdot \mathbf{x}^\top \boldsymbol{\Omega}^{-1/2} \mathbf{P}^\top \mathbf{P} \boldsymbol{\Omega}^{1/2} \boldsymbol{\Omega}^{1/2} \mathbf{P}^\top \mathbf{P} \boldsymbol{\Omega}^{-1/2} \mathbf{x} \\
 &= \chi_{n,\eta}^2 \cdot \mathbf{x}^\top \mathbf{x} \leq \chi_{n,\eta}^2,
 \end{aligned}$$

where $\mathbf{P}^\top \mathbf{P} = \mathbf{I}$ by definition and the last line follows since we know $\mathbf{x}$ is sampled from the p -ball and therefore $\mathbf{x}^\top \mathbf{x} \leq 1$. Equation (SM3.2) justifies that the $\mathbf{w}$ samples lie within the desired ellipsoid, and their uniform distribution follows from the uniform distribution of $\mathbf{x}$ in the p -ball since the distribution is preserved under a linear transformation. Finally, we accept $\mathbf{w}$ if $\mathbf{w} \in \mathcal{X}$ and reject if $\mathbf{w} \notin \mathcal{X}$. This procedure for sampling at random from $\mathcal{B}_\eta$ is summarized in Algorithm SM3.1.

Algorithm SM3.1 VGS Sampler

Input: $p \in \mathbb{N}$, $M \in \mathbb{N}$, $\mathbf{P}$, $\boldsymbol{\Omega}^{-1/2}$, $\hat{\mathbf{x}}$, $\chi_{n,\eta}^2$.

- 1: Define $\mathcal{S} := \{\}$ to be the initialized set in which the sampled points are to be placed.
- 2: **for** $k = 1, 2, \dots, M$ **do**
- 3: **Sample from the p -ball using VGS:** $\mathbf{x}_k := \mathbf{z}_{1:p} / \|\mathbf{z}\|_2$, where $\mathbf{z} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_{p+2})$ and $\mathbf{z}_{1:p}$ denotes taking the first through p -th indices (inclusive).
- 4: **Transform VGS output:** $\mathbf{w}_k := \sqrt{\chi_{n,\eta}^2} \mathbf{P} \boldsymbol{\Omega}^{-1/2} \mathbf{x}_k$.
- 5: **Translate $\mathbf{w}_k$ by the GLS estimator:** $\mathbf{w}_k \leftarrow \mathbf{w}_k + \hat{\mathbf{x}}$.
- 6: **Accept-Reject to incorporate constraints:** If $\mathbf{w}_k \in \mathcal{X}$, then add $\mathcal{S} \leftarrow \mathcal{S} \cup \{\mathbf{w}_k\}$, else start loop iteration k again.
- 7: **end for**

Output: $\mathcal{S}$ containing uniformly sampled points over $\mathcal{B}_\eta$ (as defined in (SM3.1)).

Generating samples using the VGS Sampler is efficient, but its feasibility diminishes in high dimensions because of the accept-reject step. To illustrate this point, consider the simple scenario where $\mathcal{X} = \mathbb{R}_+^p$, i.e., the non-negative orthant of $\mathbb{R}^p$ and suppose we sample from the p -ball intersected with $\mathbb{R}_+^p$ by sampling $\mathbf{x} \sim \text{VGS}(p)$ where $\mathbf{P} = \mathbf{I}$, $\hat{\mathbf{x}} = \mathbf{0}$ and

we use 1 instead of $\chi_{n,\eta}^2$. Then, $\mathbb{P}(\mathbf{x} \in \mathbb{R}_+^p) = 2^{-p}$ and the acceptance probability of each sample goes to zero exponentially in p . To make this point slightly more general, we generate data $\mathbf{y} \sim \mathcal{N}(\mathbf{x}_p^*, \mathbf{I}_p)$ where $\mathbf{x}_p^* \in \mathbb{R}_+^p$ and is a vector of ones. For a collection of dimensions $p \in [2, 30]$, we sample $\mathbf{x} \sim \text{VGS}(p)$ with $\mathbf{P} = \mathbf{\Omega} = \mathbf{I}$ and $\hat{\mathbf{x}} = \mathbf{y}$ to estimate the probability that $\mathbf{x}$ is in $\mathbb{R}_+^p$ and plot the results in the left panel of [Figure SM1](#). Since the acceptance probability decays exponentially (under 10^{-4} for 30 dimensions), it becomes clear that this algorithm is inefficient in dimensions larger than 10, since the acceptance probability quickly becomes prohibitively small in regimes where a larger sample size is even more important.

Dimension is only one of two primary complicating factors. Not only is less of a (potentially) shifted p -ball's volume in the non-negative orthant as p grows but if our data are generated via $\mathbf{y} \sim \mathcal{N}(\mathbf{K}\mathbf{x}^*, \mathbf{I}_p)$ with $\mathbf{x}^* \in \mathbb{R}_+^p$ where the condition number of $\mathbf{K}$ is large, it is possible that even less of the ellipsoid from which the VGS Sampler draws points intersects with the parameter constraint. To illustrate this point, we generate a single observation from the aforementioned model using a $\mathbf{K} \in \mathbb{R}^{40 \times 40}$ as described in [Subsection 6.3](#). The $\mathbf{x}^* \in \mathbb{R}_+^{40}$ is created in the same way as that of [Subsection 6.3.1](#). The right panel of [Figure SM1](#) shows a computed probability mass function for the number of coordinates in a VGS Sampler draw (with $\mathbf{P}$ and $\mathbf{\Omega}$ defined such that $\mathbf{K}^\top \mathbf{K} = \mathbf{P}\mathbf{\Omega}\mathbf{P}^\top$) complying with the non-negativity parameter constraints. Equivalently, the right panel of [Figure SM1](#) shows the computed probability mass function for the number of coordinates lying within the Berger–Boos set. For this particular setup, we critically note that out of 5×10^4 draws from the VGS Sampler, none of the draws had all coordinates comply with the non-negativity constraint. By contrast, for the aforementioned noise model where we more simply sample from the p -ball intersected with the non-negative orthant, the computed acceptance probability is approximately 3.35×10^{-6} , in both cases emphasizing the VGS Sampler's poor performance in high-dimensional and ill-conditioned forward model regimes.

SM3.2. Polytope sampler for general settings. In settings where the forward model is not linear and full column rank, [Algorithm SM3.1](#) fails. The ineffectiveness of this algorithm increases if the condition number of the linear forward model is large, such that most of the pre-image ellipsoid defining the Berger–Boos set lies outside of the constraint set. Although these scenarios induce particular geometric challenges, in the linear-Gaussian case with convex $\mathcal{X}$, we are still fundamentally sampling a convex set for which there is a vast literature. For example, there is a vast sampling literature for computing Bayesian posteriors in high dimensions via nested sampling [[SM18](#), [SM1](#), [SM4](#)]. In particular, nested sampling has been successfully applied in high-dimensional cosmology settings using sophisticated approaches to strategically sample the parameter space by restricting prior sampling in various ways [[SM4](#), [SM14](#)]. Although these approaches provide tools addressing a sampling setting similar to ours (i.e., the Berger–Boos set can be viewed as a portion of the parameter space defined by a cutoff on the likelihood) they are ultimately aimed at sampling from a particular distribution (i.e., the posterior), which is a stronger criterion than required here. For sampling general convex sets, simple algorithms like Hit-and-Run are available [[SM19](#), [SM10](#), [SM11](#)]. However, in the particular case considered here, more sophisticated and efficient algorithms can be devised. In particular, there exists a deep literature on random walks over polytopes such that the asymptotic stationary distribution of the walk is a uniform distribution over the polytope

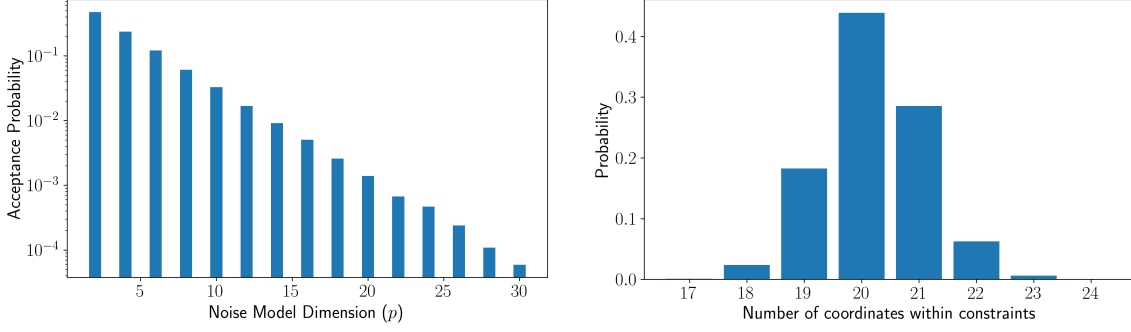


Figure SM1: Numerical illustrations of the VGS Sampler’s infeasibility in high-dimensional regimes. The **left** panel shows the computed acceptance probability of a point drawn by the VGS Sampler with data generated from a non-negatively constrained Gaussian noise model. Crucially, at only 30 dimensions, the acceptance probability is already less than 10^{-4} for this particular setup. The **right** panel shows the computed probability mass function for the number of non-negative constraint complying coordinates of a VGS sample with data generated from a non-negatively constrained linear Gaussian model in 40 dimensions with a non-identity forward model. Since this is an example using a forward model with a large condition number ($\approx 1.6 \times 10^4$), we critically note that there is empirically zero probability of generating a sample within the non-negativity constraints.

of interest [SM6, SM15, SM5]. As such, we propose to first construct a bounding polytope, $\mathcal{P}^d$ composed of d half-spaces around $\mathcal{B}_\eta$, sample C random walks within the Berger–Boos set using the Vaidya walk as described in [SM5] each starting at a point from a collection of strategically chosen locations, and then combine the parallel chains to create the final sample set. The detailed algorithm can be seen in Algorithm SM3.2.

Although MCMC algorithms are typically evaluated using trace plots on individual dimensions of the parameter space, given the high-dimensionality of the problem and the final step combining several MCMC chains in the parameter space started from different positions, it is more meaningful to evaluate this algorithm’s ability to sample fully from the functional space. Namely, we can solve for the largest and smallest values the functional can take within the Berger–Boos set as follows:

$$(SM3.3) \quad I_{BB}(\mathbf{y}) := [\mu_{BB}^l, \mu_{BB}^u] = \left[\min_{\mathbf{x} \in \mathcal{B}_\eta} \varphi(\mathbf{x}), \max_{\mathbf{x} \in \mathcal{B}_\eta} \varphi(\mathbf{x}) \right].$$

When $\varphi(\mathbf{x}) = \mathbf{h}^\top \mathbf{x}$ for some $\mathbf{h} \in \mathbb{R}^p$, $I_{BB}(\mathbf{y})$ corresponds to the SSB interval in [SM20]. Computing the sets $C_\alpha^{sl}(\mathbf{y}; \mathcal{B}_\eta)$ and $C_\alpha^{gl}(\mathbf{y}; \mathcal{B}_\eta)$ well is then contingent upon sampling functional values within $I_{BB}(\mathbf{y})$ well, since a functional value can only be included in the inverted set if the sampler has a non-zero probability of sampling arbitrarily close to it. By “well”, we informally mean that the sampled functional values range at least between the endpoints of $I_{BB}(\mathbf{y})$ and that if we partition $I_{BB}(\mathbf{y})$ into n sub-intervals, that most if not all of the sub-intervals contain at least one sample. In practice, it will often be the case that the sampled functional values can lie outside the endpoints of $I_{BB}(\mathbf{y})$ since the MCMC chains are sampling a bounding polytope of the Berger–Boos set.

Choosing the bounding polytope. Any bounded polytope of $\mathcal{B}_\eta$ is the intersection of a finite number of half-spaces defined in $\mathbb{R}^p$. The task of constructing a bounding polytope is then equivalent to choosing a collection of d hyperplanes in $\mathbb{R}^p$ to construct a set $\mathcal{P}^d := \{\mathbf{x} \in \mathbb{R}^p : \mathbf{A}\mathbf{x} \leq \mathbf{b}\}$ such that $\mathcal{B}_\eta \subseteq \mathcal{P}^d$, where $\mathbf{A} \in \mathbb{R}^{d \times p}$ and $\mathbf{b} \in \mathbb{R}^d$. Let $\mathbf{a}_i^\top \in \mathbb{R}^p$ denote the i -th row vector of $\mathbf{A}$. We compute b_i (the i -th element of $\mathbf{b}$) as

$$(SM3.4) \quad b_i = \max_{\mathbf{x} \in \mathcal{B}_\eta} \mathbf{a}_i^\top \mathbf{x}.$$

This construction ensures the necessary inclusion. We consider three approaches to pick the vectors $\mathbf{a}_i$: (i) using the constraints $\mathcal{X}$, (ii) using the known eigenvectors defining the bounded directions of the pre-image ellipsoid, and (iii) randomly. In practice, we combine these approaches to ensure that we consider only parameter settings in agreement with our physical constraints, and to tighten the bounding Berger–Boos set polytope as much as possible. There is a tradeoff with respect to the latter consideration since the mixing time and computational cost for the Vaidya walk increase with the number of hyperplanes [SM5].

To incorporate the known parameter constraints, consider the non-negativity constraint used in Section 6, i.e., $\mathbf{x} \in \mathbb{R}_+^p$. To enforce non-negativity, we set $\mathbf{a}_i = -\mathbf{e}_i$ for $i = 1, \dots, p$, where $\mathbf{e}_i$ is defined by its i -th element set to one and the rest of its elements set to zero, and $b_i = 0$. These choices produce p rows in $\mathbf{A}$ corresponding to the desired lower bounds (i.e., $x_i \geq 0$ for all i), but we can compute an additional p constraints using (SM3.4) with $\mathbf{a}_i := \mathbf{e}_i$ for $i = p + 1, \dots, 2p$. These $2p$ constraints define a hyperrectangle enclosing the Berger–Boos set in the parameter space. To incorporate polytope constraints based upon the forward model, we use the ellipsoidal definition of the pre-image as shown in (SM3.1) and eigendecomposition of $\mathbf{K}^\top \Sigma^{-1} \mathbf{K} = \mathbf{P} \mathbf{\Omega} \mathbf{P}^\top$ shown in (SM3.2). Note that this ellipsoid form is valid under the linear-Gaussian noise assumption and for all $\mathbf{K}$. In the event that $\mathbf{K}$ is not full column rank, the ellipsoid defined by Equation (SM3.1) is still defined via the pseudo-inverse of $\mathbf{K}^\top \Sigma \mathbf{K}$, but where the ellipsoid is unbounded in some directions. The column vectors of $\mathbf{P}$ corresponding to the non-zero entries on the diagonal of $\mathbf{\Omega}$ are the eigenvectors corresponding to the bounded principal axes. As such, both $\mathbf{p}_i$ and $-\mathbf{p}_i$ for $i = 1, \dots, p$ (the column vectors of $\mathbf{P}$) can be used as rows of $\mathbf{A}$ with their corresponding bounds defined by (SM3.4). Finally, to further tighten the polytope around the Berger–Boos set, we sample a multivariate Gaussian, i.e., $\mathbf{a}_i \sim \mathcal{N}(\mathbf{0}, \mathbf{I}_p)$ to include random hyperplanes. Note, it is possible that the unbounded directions of the ellipsoid defined by Equation (SM3.1) are not bounded when intersecting with the parameter constraint set, and hence there is no bounding polytope. In this case, our interval construction should be unbounded since there is not enough information to produce a bounded confidence set for the quantity of interest.

Once the components $(\mathbf{A}, \mathbf{b})$ have been defined using some or all of the above hyperplane generation strategies, the Vaidya walk can immediately be employed to perform a random walk around the polytope. The primary intuitive requirement we wish to satisfy with any sampling scheme in this context is that every region of the Berger–Boos set has a non-zero probability of being sampled. Asymptotically, the Vaidya walk samples the desired polytope uniformly at random, which satisfies a stronger requirement. In practice, the need to sample non-uniformly can arise if there are particularly meaningful parameter settings in regions that are difficult for the random walk to reach. We explore such a case in Subsection 6.2. Additionally, although

the asymptotic distribution of the Vaidya sampler is theoretically sufficient, in practice, it often has some difficulty reaching the corners of the generated polytope. Although MCMC chain mixing is typically evaluated by looking at trace plots, this diagnostic is insufficient here because the dimension is high, and the defined polytope can make different dimensions difficult to compare. Instead, we consider how well the sampler samples the functional values $\varphi(\mathbf{x})$. This motivates taking a collection of starting points between the SSB endpoints and the Chebyshev center of the polytope, as described in the following section.

Constructing the parallel chain starting points. Although the random walks defined in [SM5] asymptotically sample uniformly over $\mathcal{P}^d$, given the long and thin shape of the pre-image, running the Vaidya walk from even a “good” starting position does not consistently sample the functional space well. Instead, we use the following heuristic to construct several parallel chains that constitute a complete sample when combined. We define an even number of starting points, $C \in 2\mathbb{N}$, to roughly span the Berger–Boos set, run the Vaidya walk for M_p steps from each starting point to collect parameter settings $\{\tilde{\mathbf{x}}_1, \dots, \tilde{\mathbf{x}}_{M_p}\}$, and combine the samples from each walk, resulting in $C \times M_p$ total samples. In practice, C and M_p are chosen to yield a total of M samples as desired for Algorithm 3.1 or Algorithm 3.2. Since we have designed this process to avoid the necessity of any individual chain reaching its asymptotic distribution, we do not need a burn-in period for any chain as long as we are confident that the union of the sampled points sufficiently covers Interval (SM3.3) so that we may accept or reject any quantity of interest value in that interval. We define the collection of starting points using the line segments defined by the parameters generating the endpoints of $I_{\text{BB}}(\mathbf{y})$ and the Chebyshev center of $\mathcal{P}^d$ as defined in [SM3], and denoted by $\mathbf{x}_c$. This point is the center of the largest ball contained within $\mathcal{P}^d$ and therefore acts as a way to characterize the center of $\mathcal{P}^d$. We denote the parameter settings generating the endpoints of $I_{\text{BB}}(\mathbf{y})$ by $\hat{\mathbf{x}}^l$ for the lower endpoint and $\hat{\mathbf{x}}^u$ for the upper endpoint. We then define a uniform grid of values $\{\tau_l\}_{l=1}^{C/2}$ such that $\tau_l \in (0, 1)$ and $\tau_l < \tau_{l+1}$ for all l . For each $k \in [C]$, define $k' := \lceil k/2 \rceil$ and set $\mathbf{x}_k^{\text{start}} := \tau_{k'} \hat{\mathbf{x}}^l + (1 - \tau_{k'}) \mathbf{x}_c$ if k is odd and $\mathbf{x}_k^{\text{start}} := \tau_{k'} \hat{\mathbf{x}}^u + (1 - \tau_{k'}) \mathbf{x}_c$ if k is even. Creating starting positions along the lines connecting these endpoints and the Chebyshev center accommodates the chosen polytope while helping ensure that samples are chosen spanning the range of possible functional values over the Berger–Boos set.

The empirical performance of the Polytope sampler can be seen in Figure SM2, showing (left) a histogram of the functional values sampled and (right) a trace plot for the sampled Vaidya walks starting from points constructed as described above. These plots are generated for one observation from the 80-dimensional ill-posed inverse problem in the coverage study performed in Subsection 6.3. Critically, the histogram shows that the Polytope sampler samples the functional space well, and the trace plot shows that our starting point construction heuristic performs well in practice.

Sampling from the Berger–Boos set. With $\mathcal{P}^d$ and the starting positions defined, we construct a sample $\mathcal{S} := \{\tilde{\mathbf{x}}_1, \dots, \tilde{\mathbf{x}}_{C \times M_p}\}$. We leave the details of the Vaidya walk to the original paper [SM5], but note an essential radius tuning parameter of the algorithm that must be chosen. This radius impacts the spread of a Gaussian proposal distribution for the walk and thus affects a new proposed point’s acceptance probability. Choosing a radius that is too large results in a low acceptance rate of proposed steps, thus creating a walk that does

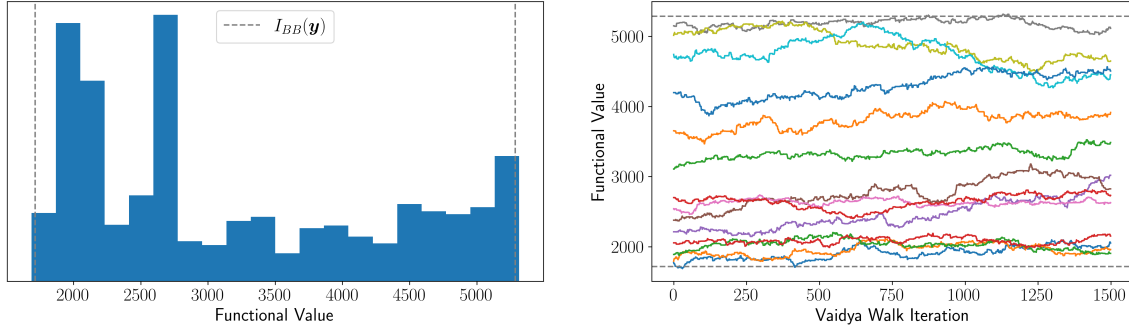


Figure SM2: Polytope sampler output for a realization of the 80-dimensional ill-posed inverse problem studied in [Subsection 6.3](#). The **left** panel contains a histogram of sampled functional values which both span and cover well the range of $I_{BB}(\mathbf{y})$ (shown by the dashed gray lines in both plots). The **right** panel contains trace plots of the 14 Vaidya walks (each indicated by a different color) which together constitute the full sample. Our heuristic for choosing starting points along the lines connecting the parameter settings generating the endpoints of $I_{BB}(\mathbf{y})$ and the Chebyshev center of the polytope provides a good initial spread of starting functional values.

not mix well. In contrast, choosing a radius that is too small results in a high acceptance rate with relatively small step sizes. In practice, we find a radius setting of 0.5 works well as it produces an acceptance probability of $\approx 33.3\%$, producing a reasonable tradeoff between taking meaningful steps and not rejecting too many steps.

SM4. Quantile regression overview. Fundamentally, given a one-dimension random variable $z \sim P_{\mathbf{x}}$ that depends on the parameter $\mathbf{x} \in \mathbb{R}^p$, we are interested in the upper γ -quantile at every parameter setting, i.e.,

$$(SM4.1) \quad \mathbb{P}_{\mathbf{x}}(z > Q_{\mathbf{x}}(1 - \gamma)) = \gamma.$$

We note that the quantile surface itself is not random, so we can use draws from the distribution $P_{\mathbf{x}}$ to estimate $Q_{\mathbf{x}}$ at a given parameter setting $\mathbf{x}$. However, as noted in [Subsection 3.2](#), performing such an estimate is not always computationally feasible, and intuitively, we might expect quantiles to vary smoothly over the parameter space, which would imply that information about a quantile at $\mathbf{x}_1$ should be related to a quantile at $\mathbf{x}_2$ if these points are close. In statistics literature, estimating $Q_{\mathbf{x}}$ is framed as estimating a quantile function conditional on known covariates and is often thought of as a generalization of estimating the conditional median [\[SM9\]](#). Just as conditional mean and conditional median estimation can be accomplished by using an appropriate loss function (sum of squares and absolute differences, respectively), estimating conditional quantiles can be accomplished by minimizing the pinball loss defined as follows:

$$(SM4.2) \quad L_{\gamma}(z, q) := \begin{cases} (1 - \gamma)(q - z), & z < q, \\ -\gamma(z - q), & z \geq q. \end{cases}$$

[\[SM7, SM21\]](#). As such, estimating the quantile surface can be framed as a risk-minimization problem, leaving only standard modeling choices to fill in for q in [\(SM4.2\)](#). Although initial

Algorithm SM3.2 Polytope sampler

Input: $M_p \in \mathbb{N}$ and $C \in 2\mathbb{N}$. $\mathbf{A} \in \mathbb{R}^{d \times p}$, $\mathbf{b} \in \mathbb{R}^d$. $\{\tau_l\}_{l=1}^{C/2}$, where $\tau_l \in (0, 1)$ and $\tau_l < \tau_{l+1}$ for all l .

- 1: Let $\mathcal{S} := \{\}$ be the set in which we store all sampled parameter settings.
- 2: **Construct Chebyshev center:** Solve for $\mathbf{x}_c$ using the following optimization.

$$\begin{aligned} & \underset{\mathbf{x}_c, r}{\text{maximize}} && r \\ & \text{subject to} && \mathbf{a}_i^\top \mathbf{x}_c + r \|\mathbf{a}_i\|_2 \leq b_i, \quad i = 1, \dots, d. \end{aligned}$$

- 3: **Compute φ extremes of Berger–Boos set:** Extreme points are computed with respect to the functional of interest:

$$\begin{aligned} \widehat{\mathbf{x}}^l &:= \operatorname{argmin}_{\mathbf{x} \in \mathcal{B}_\eta} \varphi(\mathbf{x}) \\ \widehat{\mathbf{x}}^u &:= \operatorname{argmax}_{\mathbf{x} \in \mathcal{B}_\eta} \varphi(\mathbf{x}) \end{aligned} \quad \text{(SM3.5)}$$

- 4: **for** $k = 1, 2, \dots, C$ **do**
- 5: **Construct starting point:** Define $k' := \lceil k/2 \rceil$. Define $\mathbf{x}_k^{\text{start}} = \tau_{k'} \widehat{\mathbf{x}}^l + (1 - \tau_{k'}) \mathbf{x}_c$ if k is odd and $\mathbf{x}_k^{\text{start}} = \tau_{k'} \widehat{\mathbf{x}}^u + (1 - \tau_{k'}) \mathbf{x}_c$ if k is even.
- 6: **Run the Vaidya walk for M_p steps:** Collect samples $\{\tilde{\mathbf{x}}_1, \dots, \tilde{\mathbf{x}}_{M_p}\}$ and add them to $\mathcal{S}$.
- 7: **end for**

Output: $\mathcal{S}$ containing sampled points over $\mathcal{B}_\eta$.

efforts were focused on linear parametric quantile regressors [SM8], in recent years, modeling efforts have focused on nonparametric varieties. [SM13] adapted random forests to quantile regression. [SM22] leveraged Reproducing Kernel Hilbert Spaces to construct smooth quantile regressors. Closer to our application, [SM12] used neural networks to optimize the pinball loss to learn the quantile surface for their application.

SM5. Additional details and illustrations in Section 6.

SM5.1. Importance-like sampler for the three-dimensional example in Subsection 6.2.

Since the parameter settings with quantiles meaningfully larger than $\chi_{1,\alpha}^2$ are located close to the constraint boundary, a sampling challenge is presented. Using the samplers as described in Subsection 5.1 results in under-sampling of this large-quantile region since both samplers provide uniform random samples over the Berger–Boos set. Algorithm SM5.1 presents a modified version of Algorithm 3.1, an importance-like sampler to increase the probability mass of samples close to the constraint boundary. Note, we say “importance-like” because we do not provide any theoretical guarantee regarding this sampler’s ability to produce draws from a particular target distribution. We tailored Algorithm SM5.1 to settings with a non-negativity constraint and hand-tuned the length scale parameter to the particular three-dimensional

example in [Subsection 6.2](#).

Algorithm SM5.1 Importance-like sampler for three-dimensional constrained Gaussian

Input: Number of samples: $M \in \mathbb{N}$, inverse length scale: γ_p , and order of norm: $q \in (0, 1)$.

```

1: Instantiate a list  $\mathcal{S}$  of length  $M$  to store sampled points.
2: while  $|\mathcal{S}| < M$  do
3:   Draw  $M - |\mathcal{S}|$  realizations from Algorithm SM3.2:  $\tilde{\mathbf{x}}_1, \dots, \tilde{\mathbf{x}}_{M-|\mathcal{S}|}$ 
4:   for  $k = 1, \dots, M - |\mathcal{S}|$ : do
5:     Compute the probability of accepting the  $k$ -th draw:  $p_k := \exp(-\gamma_p \|\tilde{\mathbf{x}}_k\|_q)$ .
6:     Draw  $z_k \sim \text{Bernoulli}(p_k)$ .
7:     if  $z_k = 1$  then
8:        $\mathcal{S}[k] \leftarrow \tilde{\mathbf{x}}_k$ 
9:     end if
10:  end for
11: end while

```

Output: Sampled parameters in Berger–Boos set $\mathcal{S}$.

The key to [Algorithm SM5.1](#) is the additional accept/reject step where the k -th sample is accepted with probability p_k . Additionally, by setting $q \in (0, 1)$, we prioritize retaining samples closer to the non-negativity boundary. This effect can be seen in [Figure SM3](#), where the vast majority of samples are found closer to the constraint boundary. The ability of [Algorithm SM5.1](#) to better sample the high-quantile regions of the Berger–Boos set can be seen in [Figure SM4](#). In particular, for the functional values $\mu \in (-4, -2)$, the importance-like sampler is substantially more effective than the Polytope sampler at finding parameter settings with γ -quantiles greater than 1.1.

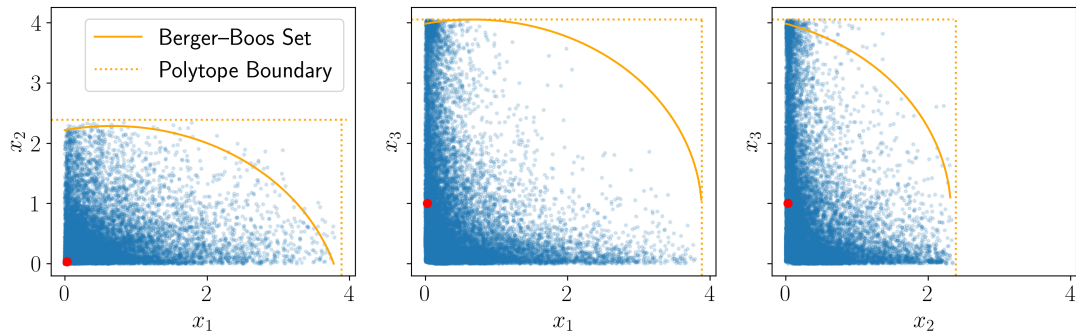


Figure SM3: When sampling points using [Algorithm SM5.1](#), the vast majority of samples are found closer to the non-negativity constraint boundary. The true parameter setting is shown by the red point, while the parameter settings sampled by [Algorithm SM5.1](#) are shown by the blue points. This sampling prioritization helps adequately sample the regions of the Berger-Boos set where the quantile surface is larger than $\chi^2_{1,\alpha}$. Furthermore, the vast majority of sampled points lie within the Berger-Boos set, with some lying outside within the bounding polytope.

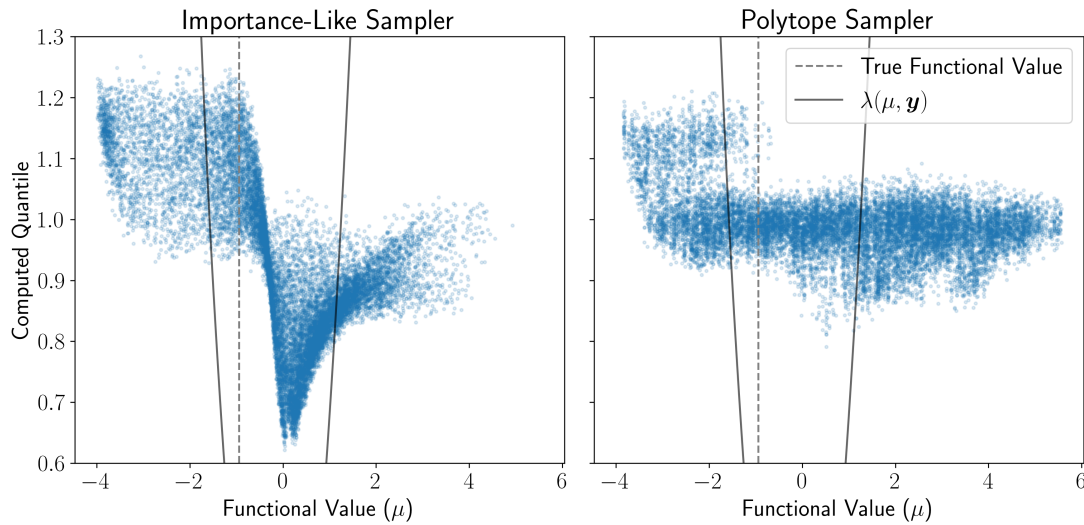


Figure SM4: The importance-like sampler described by [Algorithm SM5.1](#) is more effective than the Polytope sampler described by [Algorithm SM3.2](#) at sampling parameter settings with γ -quantile greater than 1.1. Each parameter setting sampled by [Algorithm SM5.1](#) is shown by a blue point. This improved ability helps ensure the coverage guarantee shown in the left panel of [Figure 4](#).

REFERENCES

- [1] G. ASHTON ET AL., *Nested sampling for physical scientists*, Nature Reviews Methods Primers, 2 (2022).
- [2] R. L. BERGER AND D. D. BOOS, *P values maximized over a confidence set for the nuisance parameter*, Journal of the American Statistical Association, 89 (1994), pp. 1012–1016.
- [3] S. BOYD AND L. VANDENBERGHE, *Convex Optimization*, Cambridge University Press, 2004.
- [4] J. BUCHNER, *Nested sampling methods*, arXiv preprint arXiv:2101.09675v4, (2023).
- [5] Y. CHEN, R. DWIVEDI, M. J. WAINWRIGHT, AND B. YU, *Fast mcmc sampling algorithms on polytopes*, Journal of Machine Learning Research, 19 (2018), pp. 1–86.
- [6] R. KANNAN AND H. NARAYANAN, *Random walks on polytopes and an affine interior point method for linear programming*, Mathematics of Operations Research, 37 (2012), pp. 1 – 20.
- [7] R. KOENKER, *Quantile regression*, vol. 38, Cambridge university press, 2005.
- [8] R. KOENKER AND G. BASSETT JR, *Regression quantiles*, Econometrica: journal of the Econometric Society, (1978), pp. 33–50.
- [9] R. KOENKER AND K. F. HALLOCK, *Quantile regression*, Journal of Economic Perspectives, 15 (2001), pp. 143–156.
- [10] L. LOVASZ, *Hit-and-run mixes fast*, Mathematical Programming, 86 (1999), pp. 443 – 461.
- [11] L. LOVASZ AND S. VEMPALA, *Hit-and-run from a corner*, SIAM Journal on Computing, 35 (2006), pp. 985 – 1005.
- [12] L. MASSERANO, T. DORIGO, R. IZBICKI, M. KUUSELA, AND A. LEE, *Simulation-based inference with waldo: Confidence regions by leveraging prediction algorithms or posterior estimators for inverse problems*, AISTATS, (2023).
- [13] N. MEINSHAUSEN, *Quantile regression forests*, Journal of Machine Learning Research, 7 (2006), pp. 983–999.
- [14] N. A. MONTEL, J. ALVEY, AND C. WENIGER, *Scalable inference with autoregressive neural ratio estimation*, arXiv preprint arXiv:2308.08597v1, (2023).
- [15] H. NARAYANAN, *Randomized interior point methods for sampling and optimization*, The Annals of Applied Probability, 26 (2016), pp. 597–641.
- [16] E. A. OK, *Real Analysis with Economic Applications*, Econometric Society Monographs, Princeton University Press, Princeton, NJ, 1st ed., 2007, <https://press.princeton.edu/books/hardcover/9780691129495/real-analysis-with-economic-applications>. Includes bibliographical references and index.
- [17] B. W. RUST AND W. R. BURRUS, *Mathematical Programming and the Numerical Solution of Linear Equations*, American Elsevier, 1972.
- [18] J. SKILLING, *Nested sampling*, AIP Conference Proceedings, 735 (2004), p. 295.
- [19] R. L. SMITH, *Efficient monte carlo procedures for generating points uniformly distributed over bounded regions*, Operations Research, 32 (1984), pp. 1296–1308.
- [20] M. STANLEY, P. PATIL, AND M. KUUSELA, *Uncertainty quantification for wide-bin unfolding: one-at-a-time strict bounds and prior-optimized confidence intervals*, Journal of Instrumentation, 17 (2022), <https://doi.org/https://doi.org/10.1088/1748-0221/17/10/P10013>.
- [21] I. STEINWART AND A. CHRISTMANN, *Estimating conditional quantiles with the help of the pinball loss*, Bernoulli, 117 (2011), pp. 211–225.
- [22] I. TAKEUCHI, Q. V. LE, T. D. SEARS, AND A. J. SMOLA, *Nonparametric quantile estimation*, Journal of Machine Learning, 7 (2006), pp. 1231–1264.
- [23] A. R. VOELKER, J. GOSMANN, AND T. C. STEWART, *Efficiently sampling vectors and coordinates from the n-sphere and n-ball*, tech. report, Centre for Theoretical Neuroscience, 2017.