

PRIOR DISTRIBUTIONS FOR GAUSSIAN MODELS UNDER DIFFERENTIAL PRIVACY WITH BOUNDS ON DATA VALUES

ZEKI KAZAN AND JEROME P. REITER

Department of Statistical Science, Duke University
e-mail address: zekicankazan@gmail.com

Department of Statistical Science, Duke University
e-mail address: jreiter@duke.edu

ABSTRACT. The Gaussian distribution is often used to model univariate, numerical data, even in contexts where individuals' data values are bounded, e.g., test scores and biological measurements. We consider Bayesian inference for the parameters of Gaussian models when (i) analysts seek a post-processing algorithm that takes noisy sufficient statistics as inputs and produces a posterior distribution of the parameters as output, and (ii) the privacy mechanism enforces bounds on individual data values as a condition for satisfying pure differential privacy. Our primary contributions focus on the specification of prior distributions for model parameters. Specifically, we illustrate that failing to account for the bounds on data values when using informative prior distributions can result in posterior inferences with undesirable properties. Further, we show theoretically that the usual default priors used in the nonprivate setting do not lead to proper posterior distributions, and we suggest a default prior distribution that does. We show empirically that one can reduce mean squared errors and avoid inferences with infeasible support by incorporating the boundedness assumptions into the specification of prior distributions. We present methods for inferring about the mean and variance of a univariate Gaussian model and the coefficients and variance of a linear regression. As we discuss, the general strategies used for the Gaussian models can be adapted for other models when enforcing bounds on the data values to ensure differential privacy.

1. INTRODUCTION

Gaussian models are among the most commonly used in applied statistics. They are used, for example, to obtain inferences from sampled data for population-level means and variances, as well as for population-level coefficients of linear regression models. Although the Gaussian distribution has infinite support, it is also frequently used for data that are bounded. For example, Gaussian models are typical for quantities like proportions and standardized test scores, which have known ranges. They are also used for quantities like health measurements (e.g., blood pressures, birth weights) and stock market returns, which can be viewed as bounded.

Key words and phrases: Bayesian, Confidentiality, Posterior, Regression, Sampling.

We consider contexts in which a data analyst makes inferences about the parameters of a Gaussian model under differential privacy (DP) (Dwork et al., 2006). We focus specifically on settings where the analyst has access to differentially private versions of the empirical mean and variance, which are sufficient statistics for the univariate Gaussian distribution. Given these noisy statistics, the analyst intends to apply a post-processing algorithm that produces a posterior distribution of the parameters; that is, the analyst intends to use Bayesian inference. We suppose the DP mechanism involves bounding the data values to ensure finite global sensitivity, for example, by imposing the assumption that confidential values lie within an interval $[a, b]$. Enforcing bounds is a common approach to constructing DP algorithms (Bernstein and Sheldon, 2019; Du et al., 2020; Sheffet, 2017; Wang, 2018; Zhang et al., 2012; Kamath et al., 2025).

As with any Bayesian inference, the analyst must specify a prior distribution for the Gaussian model parameters. To the best of our knowledge, existing applications in the literature generally presume the analyst uses an informative prior distribution. In practice, however, the analyst may not have sufficient prior knowledge to set a reasonable informative prior distribution. Additionally, applications in the literature typically rely on prior distributions with unbounded support. This can be in conflict with the constraints; for example, it could be nonsensical for the prior distribution to have mass on values of the mean parameter outside $[a, b]$ when all data values must lie in that region.

In this article, we address incorporating constraints on the data values in Bayesian inference for Gaussian models under DP, focusing on implications for prior specification. Our primary contributions include:

- We illustrate that failing to account for the bounds on data values when using informative prior distributions can result in posterior inferences with undesirable properties. This finding accords with research showing that Bayesian inference under DP can be sensitive to the choice of prior distribution (e.g., Minami et al., 2016; Zhang et al., 2016).
- We develop an improved Gibbs sampler for Bayesian inference for univariate Gaussian models, given noisy sufficient statistics. The sampler uses the exact forms of all full conditional distributions, avoiding approximations that are employed in other Bayesian post-processing inferential approaches in the literature.
- We show theoretically that the improper priors typically used for univariate Gaussian models in the nonprivate setting do not lead to proper posterior distributions, and we suggest a default prior distribution that does.
- We propose a straightforward method to incorporate bounds by constraining the support of the distribution for the (unknown) values of the confidential sufficient statistics, as well as the prior distribution on the parameters, to accord with the bounds on the data values. We show empirically that applying the method avoids inferences with infeasible support and can reduce mean squared errors relative to methods that disregard the constraints.

We emphasize that our goal is to examine the impact of prior specification on Bayesian inference with differential privacy, not to find Bayesian methods with well-calibrated frequentist properties. Bayesian posterior intervals are not guaranteed to have nominal coverage rates, as the prior distribution can and does influence the interval. Indeed, this is the case for our method when using the default prior distributions; this is evident in the simulation results presented later.

The remainder of this article is organized as follows. In Section 2, we briefly review DP and some work related to Bayesian inference with DP. Section 3 and Section 4 concern

Bayesian inference for the mean and variance of a univariate Gaussian model using a measurement error framework, that is, we model the noisy statistics given the values of the confidential statistics (which are unknown to the analyst) and model the (unknown) values of the confidential statistics given parameters of the Gaussian model. Specifically, Section 3 illustrates inferences when using informative prior distributions, including the Gibbs sampler that uses exact full conditional distributions. Section 4 explains how the usual noninformative prior distributions used in Bayesian inference in the nonprivate setting fail to result in proper posterior distributions, and presents our suggested default prior distribution that respects the constraints on the data values. This section also includes simulation studies that examine repeated sampling properties to compare Bayesian inferences with the default prior distribution to those with a prior distribution that disregards the constraints. In Section 5, we turn from the univariate Gaussian distribution to Bayesian inference for the parameters of a simple linear regression under DP. Here, we use measurement error models for post-processing the noisy sufficient statistics, accounting for the constraints on the sufficient statistics implied by the bounds on the variables in the model. Finally, in Section 6, we summarize the findings and discuss future research directions. In particular, we suggest how incorporating constraints on the sufficient statistics and prior support could be used for Bayesian inference for the parameters of other models. All code can be found at https://github.com/zekicankazan/dp_priors.

2. BACKGROUND

We develop methods for use when the Laplace mechanism is used to satisfy pure DP. For convenience, we briefly define each here. We expect other variants of DP and other mechanisms to raise similar issues for prior specification. We then review some related work on Bayesian inference and DP.

2.1. Differential Privacy. We begin with the formal definition of DP, introduced by [Dwork et al. \(2006\)](#).

Definition 2.1 (Differential Privacy). A mechanism $\mathcal{M}$ satisfies ε -differential privacy if for all data sets x, x' that differ in only one row and all $S \subseteq \text{Range}(\mathcal{M})$,

$$P[\mathcal{M}(x) \in S] \leq e^\varepsilon P[\mathcal{M}(x') \in S]. \quad (2.1)$$

DP has many desirable properties ([Dwork et al., 2006](#)). Here, we make use of post-processing, which ensures that functions of DP statistics maintain the privacy guarantee.

Theorem 2.2 (Post-processing). *For any data set x , if $\mathcal{M}(x)$ satisfies ε -DP and g is not a function of x , then $g(\mathcal{M}(x))$ satisfies ε -DP.*

We also make use of composition, which provides a method for composing the privacy guarantees from multiple DP releases.

Theorem 2.3 (Composition). *For any data set x , if $m_1 = \mathcal{M}_1(x)$ satisfies ε_1 -differential privacy and $m_2 = \mathcal{M}_2(x, m_1)$ satisfies ε_2 -differential privacy, then $\mathcal{M}(x) = (m_1, m_2)$ satisfies $(\varepsilon_1 + \varepsilon_2)$ -differential privacy.*

The (global) sensitivity of a function is defined as follows.

Definition 2.4 (Sensitivity). For any function f applied to any data set x , the sensitivity is defined as $\Delta f = \max_{x, x'} |f(x) - f(x')|$ for all data sets x, x' that differ in only one row.

The sensitivity of a function f is used in the Laplace mechanism, which is a common method for achieving DP.

Theorem 2.5 (Laplace Mechanism). *If a function $f(x)$ has sensitivity Δf , then the mechanism $\mathcal{M}(x) = \text{Lap}(f(x), \Delta f/\varepsilon)$ satisfies ε -DP.*

2.2. Related Work. Bayesian methods are natural for statistical inference after a DP release. By sampling from posterior distributions, analysts can perform estimation and prediction tasks automatically with “built-in” uncertainty quantification. As examples, [Bernstein and Sheldon \(2018\)](#) and [Ju et al. \(2022\)](#) use Gibbs sampling to estimate differentially private posterior distributions, the former for exponential families and the latter for likelihoods satisfying a record additivity condition. [Gong \(2022\)](#) uses approximate Bayesian computation for settings where a perturbation mechanism is applied. Other examples include Bayesian estimation of subset proportions ([Li and Reiter, 2022](#)), linear regression models ([Bernstein and Sheldon, 2019](#)), and generalized linear models ([Kulkarni et al., 2021](#)). These methods employ variations of the following scheme: repeatedly impute plausible values of the nonprivate statistics based on their private counterparts, and obtain samples of the parameters of interest via a nonprivate Bayesian analysis.

DP is related to Bayesian methods more generally. Releasing a sample from a posterior distribution can be used to achieve DP ([Dimitrakakis et al., 2017](#); [Geumlek et al., 2017](#); [Jewson et al., 2024](#); [Minami et al., 2016](#); [Wang et al., 2015](#); [Zhang and Zhang, 2023](#); [Zhang et al., 2016](#)). Another line of work examines Bayesian semantics for DP and their relationship to disclosure risk assessment ([Dwork et al., 2006](#); [Kasiviswanathan and Smith, 2014](#); [Kazan and Reiter, 2024, 2025](#); [Kifer et al., 2022](#); [Kifer and Machanavajjhala, 2014](#); [McClure and Reiter, 2012](#)).

Finally, we note that Gaussian models have been studied in DP contexts by many authors, e.g., [Biswas et al. \(2020\)](#); [Kulesza et al. \(2024\)](#); [D’Orazio et al. \(2015\)](#); [Du et al. \(2020\)](#); [Ferrando et al. \(2022\)](#); [Karwa and Vadhan \(2018\)](#); [Couch et al. \(2019\)](#); [Kazan et al. \(2023\)](#); [Peña and Barrientos \(2025\)](#). To the best of our knowledge, these authors have not addressed the issues in prior specification for Bayesian inference that we consider here.

3. MEASUREMENT ERROR MODELING AND CONSTRAINTS FOR UNIVARIATE GAUSSIAN

Let $Y_1, \dots, Y_n$ be confidential, scalar data values for n individuals. The Y_i are contained in some interval $[a, b]$ set by the party responsible for implementing DP, who we henceforth refer to as the data curator. The data curator sets the interval to accord with scientific understanding of the phenomenon under study and, to maintain the DP guarantee, independently of the realized data values. As examples, test scores may be bounded between $[a, b] = [0, 100]$ and IQ scores between $[a, b] = [40, 250]$. This interval is made public by the data curator.

The data curator releases private versions of the sample mean and variance, i.e., $\bar{Y} = \sum_{i=1}^n Y_i/n$ and $S^2 = \sum_{i=1}^n (Y_i - \bar{Y})^2/(n-1)$. As shown in [Du et al. \(2020\)](#), $\bar{Y}$ and S^2 have sensitivities $(b-a)/n$ and $(b-a)^2/n$, respectively. To achieve ε -DP, we presume the

data curator releases $\bar{Y}^* \sim \text{Lap}(\bar{Y}, (b-a)/(\varepsilon_1 n))$ and $S^{2*} \sim \text{Lap}(S^2, (b-a)^2/(\varepsilon_2 n))$, where $\varepsilon_1 + \varepsilon_2 = \varepsilon$.

We make the simplifying assumptions $a = 0$ and $b = 1$. We can do so without loss of generality, as described in Observation 3.1, which is proved in Appendix A.

Observation 3.1. Let $Y_1, \dots, Y_n \in [a, b]$ and let $\tilde{Y}_i = (Y_i - a)/(b - a) \in [0, 1]$. Let $\bar{Y}$ and S^2 be the sample mean and variance for $\{Y_i\}$ and let $\tilde{\bar{Y}}$ and $\tilde{S}^2$ be the sample mean and variance for $\{\tilde{Y}_i\}$. Suppose each statistic is released via the Laplace mechanism under ε -DP and denote the DP statistics $\bar{Y}^*$, S^{2*} , $\tilde{\bar{Y}}^*$, and $\tilde{S}^{2*}$. Then, $\bar{Y}^* \stackrel{d}{=} (b-a)\tilde{\bar{Y}}^* + a$ and $S^{2*} \stackrel{d}{=} (b-a)^2\tilde{S}^{2*}$.

A data analyst with access to only $\bar{Y}^*$ and S^{2*} , as well as the particulars of the DP mechanism used to create them, considers these released summary statistics as noisy versions of $\bar{Y}$ and S^2 , which themselves are estimates of underlying population parameters μ and σ^2 . To make inferences for μ and σ^2 , we presume the analyst posits a model for the underlying confidential data, namely $Y_i \sim \mathcal{N}(\mu, \sigma^2)$ independently for all i . While theoretically this model has support outside the interval $[a, b]$, in practice it accurately describes data distributions when the observed values are symmetric with an empirical range not near the bounds.

We consider an analyst who follows the Bayesian paradigm to obtain posterior inferences about (μ, σ^2) given $(\bar{Y}^*, S^{2*})$. We consider first the case of an analyst who incorporates prior beliefs into their analysis via the convenient, informative prior distribution,

$$\sigma^2 \sim \text{IG}(\nu_0/2, \nu_0\sigma_0^2/2), \mu \mid \sigma^2 \sim \mathcal{N}(\mu_0, \sigma^2/\kappa_0) \quad (3.1)$$

for analyst-specified hyperparameters $(\nu_0, \kappa_0, \mu_0, \sigma_0^2)$. Here, **IG** denotes the inverse-gamma distribution. The prior distribution in (3.1) is conjugate for the normal model, which facilitates computation.

3.1. Inference when Disregarding the Constraints. Using (3.1), analysts who disregard the constraints in post-processing inference can sample from the posterior distribution $p(\mu, \sigma^2 \mid \bar{Y}^*, S^{2*})$ using Monte Carlo methods. One approach is the Gibbs sampler of [Bernstein and Sheldon \(2018\)](#) for exponential families. In their sampler, the full conditional distribution of S^2 is approximated as Gaussian. Since S^2 must be positive, this approximation can lead to unreliable results, particularly when σ^2 is near zero or n is small. The approximation allows the sampler to have computational complexity $\mathcal{O}(1)$ per Gibbs iteration.

We now present our first contribution: when disregarding constraints, one can modify the sampler of [Bernstein and Sheldon \(2018\)](#) to draw from the exact conditional distribution of S^2 . This avoids the potential shortcomings of a Gaussian approximation for an estimated variance. Specifically, when $\varepsilon_2 < 2$, the full conditional for S^2 can be shown to be a distribution we call a truncated gamma mixture (TGM).¹

Definition 3.2. $X \sim \text{TGM}(\alpha, \beta, \lambda, \tau)$ for $\alpha > 0$, $\tau \in \mathbb{R}$, and $\beta > \lambda \geq 0$ if its distribution is as follows. When $\tau \leq 0$, $X \sim \text{Gamma}(\alpha, \beta + \lambda)$ and when $\tau > 0$, it has probability density

¹For the distribution of S^2 to be a TGM, we require $\sigma^2 < (n-1)/(2n\varepsilon_2)$. We show in Observation 3.4 that in this setting we can guarantee $\sigma^2 \leq 1/4$, which implies $\sigma^2 < (n-1)/(2n\varepsilon_2)$ automatically when $\varepsilon_2 < 2$. Operationally, we sample from a truncated full conditional for σ^2 within the Gibbs sampler, where $\sigma^2 < (n-1)/(2n\varepsilon_2)$.

function

$$p(x) = \begin{cases} \pi_1 \frac{(\beta-\lambda)^\alpha}{\gamma(\alpha, (\beta-\lambda)\tau)} x^{\alpha-1} e^{-(\beta-\lambda)x}, & \text{if } x \leq \tau; \\ \pi_2 \frac{(\beta+\lambda)^\alpha}{\Gamma(\alpha, (\beta+\lambda)\tau)} x^{\alpha-1} e^{-(\beta+\lambda)x} & \text{if } x > \tau, \end{cases} \quad (3.2)$$

where

$$\begin{aligned} \pi_1 &= \frac{e^{-\lambda\tau} \frac{\gamma(\alpha, (\beta-\lambda)\tau)}{(\beta-\lambda)^\alpha}}{e^{-\lambda\tau} \frac{\gamma(\alpha, (\beta-\lambda)\tau)}{(\beta-\lambda)^\alpha} + e^{\lambda\tau} \frac{\Gamma(\alpha, (\beta+\lambda)\tau)}{(\beta+\lambda)^\alpha}}, \\ \pi_2 &= \frac{e^{\lambda\tau} \frac{\Gamma(\alpha, (\beta+\lambda)\tau)}{(\beta+\lambda)^\alpha}}{e^{-\lambda\tau} \frac{\gamma(\alpha, (\beta-\lambda)\tau)}{(\beta-\lambda)^\alpha} + e^{\lambda\tau} \frac{\Gamma(\alpha, (\beta+\lambda)\tau)}{(\beta+\lambda)^\alpha}}. \end{aligned} \quad (3.3)$$

Here, $\gamma(\alpha, x)$ and $\Gamma(\alpha, x)$ are the lower and upper incomplete gamma functions. We provide an algorithm for sampling from the TGM in Appendix B.

In this modified Gibbs sampler, the full conditionals for $(\mu, \sigma^2, \bar{Y})$ correspond to those of Bernstein and Sheldon (2018) applied to the univariate Gaussian setting. That is, the full conditionals for μ and σ^2 have the same form as in the nonprivate setting. The full conditional for $\bar{Y}$ is based on the Bayesian LASSO (Park and Casella, 2008), whereby the fact that $\bar{Y}^* \sim \text{Lap}(\bar{Y}, 1/(\varepsilon_1 n))$ is equal in distribution to $\bar{Y}^* \sim \mathcal{N}(\bar{Y}, \omega^2)$ for $\omega^2 \sim \text{Exp}(\varepsilon_1^2 n^2/2)$ implies that the full conditional for $\bar{Y}$ is Gaussian and the full conditional for $1/\omega^2$ is inverse-Gaussian. Parameters of the full conditionals and explicit algorithms for our modified Gibbs sampler are provided in Appendix C.

3.2. Inference when Enforcing the Constraints. We now show that analysts who incorporate the fact that the data values lie within some $[a, b]$ can reduce the uncertainty in estimates of parameters of interest “for free.” They need not make additional assumptions; they merely determine how conditions on the data values constrain their model.

We incorporate the constraints on the data values by turning them into implied constraints on all unknown parameters, including the unobserved sufficient statistics from the confidential data, which we then use as part of the modeling and sampling. This estimation method can be generalized beyond the univariate Gaussian model, as we show in Section 5 and discuss in Section 6. For the univariate Gaussian model specifically, we start with the computationally efficient, modified sampler of Bernstein and Sheldon (2018) that uses the exact full conditional for S^2 described in Section 3.1. We then add constraints to this sampler, as we describe in Section 3.2.1.

We note that Bernstein and Sheldon (2018) propose a Gibbs sampler that utilizes constraints for univariate data modeled via an exponential family. Their constrained sampler still requires the Gaussian approximation for S^2 and is substantially more computationally intensive than their unconstrained one.

3.2.1. Constraints for the Univariate Gaussian Model. The constraint that $Y_i \in [0, 1]$ affects the model in two ways. Firstly, the bound constrains the parameters μ and σ^2 , since bounded random variables have bounded moments. Secondly, the bound constrains the sufficient statistics $\bar{Y}$ and S^2 , since samples of bounded random variables have bounded means and variances. We now establish these bounds in a manner that can be incorporated into the Gibbs sampler from Section 3.1. See Appendix D for proofs.

First, Observation 3.3 establishes conditional bounds for $\mu \mid \sigma^2$ and $\sigma^2 \mid \mu$, which follow from monotonicity of expectation.

Observation 3.3. Let $Y_i \in [0, 1]$ have moments $E[Y_i] = \mu$ and $V[Y_i] = \sigma^2$. It follows that $\sigma^2 \in [0, \mu(1 - \mu)]$ and $\mu \in [1/2 - \sqrt{1/4 - \sigma^2}, 1/2 + \sqrt{1/4 - \sigma^2}]$.

Applying Observation 3.3 yields the following.

Observation 3.4. If $Y_i \in [0, 1]$, then $V[Y_i] = \sigma^2 \leq 1/4$.

Observation 3.5 establishes conditional bounds for $\bar{Y} \mid S^2$ and $S^2 \mid \bar{Y}$.

Observation 3.5. Let $Y_1, \dots, Y_n$ be such that each $Y_i \in [0, 1]$. Then $S^2 \in [0, n/(n-1)\bar{Y}(1 - \bar{Y})]$ and $\bar{Y} \in [1/2 - \sqrt{1/4 - (n-1)/n \cdot S^2}, 1/2 + \sqrt{1/4 - (n-1)/n \cdot S^2}]$.

Given these observations, we can replace the unbounded normal-inverse gamma prior from (3.1) with a prior of the same form truncated to be within the feasible region, which we denote A and define as

$$A = \{(\mu, \sigma^2) : \mu \in [0, 1] \text{ and } \sigma^2 \in (0, \mu(1 - \mu))\}. \quad (3.4)$$

This yields a prior distribution of the form,

$$\sigma^2 \sim \text{IG}(\nu_0/2, \nu_0\sigma_0^2/2), \quad \mu \mid \sigma^2 \sim \mathcal{N}(\mu_0, \sigma^2/\kappa_0) \quad \text{s.t.} \quad (\mu, \sigma^2) \in A. \quad (3.5)$$

The resulting posterior distribution has a similar form as in Section 3.1, although the addition of constraints can result in dramatic changes to point and interval estimates, as we now demonstrate. The full conditionals under these constraints are described in Appendix C.

3.2.2. Example: The Blood Lead Dataset. We now demonstrate the effect of incorporating these constraints using genuine data. Computations are performed on the $[0, 1]$ scale, but all quantities are converted back to the original scale for clarity of presentation. Interval estimates are those with highest posterior density (HPD). The following example, adapted from an introductory statistics textbook (Diez et al., 2012, Section 7.1), is representative of common statistical analyses for numerical, scalar data.

Example 3.6. Researchers sampled the blood lead levels of $n = 43$ policemen assigned to outdoor work in Egypt to examine the effect of leaded gasoline on exposed individuals (Mortada et al., 2001). The sample mean and variance are $\bar{Y} = 32.08 \mu\text{g/dL}$ and $S^2 = 16.98^2 \mu\text{g}^2/\text{dL}^2$, respectively. An expert advises that blood lead levels in this region are reasonably bounded above by $100 \mu\text{g/dL}$. Using these bounds, the researchers use the Laplace mechanism with $\varepsilon_1 = \varepsilon_2 = 0.25$ to release noisy statistics $\bar{Y}^* = 34.30 \mu\text{g/dL}$ and $S^{2*} = 47.16^2 \mu\text{g}^2/\text{dL}^2$. A secondary data analyst seeks inferences for μ and σ^2 , the average and variance of blood lead levels of all policemen assigned to outdoor work in Egypt. She believes blood lead levels are reasonably approximated by a Gaussian distribution. To reflect prior knowledge, she uses (3.5) with $(\mu_0 = 12.5, \sigma_0^2 = 3.8^2, \kappa_0 = 1, \nu_0 = 1)$, roughly representing the equivalent information of one “prior data point.”²

Figure 1 presents the joint posterior distribution for (μ, σ^2) . It also displays the posterior distribution when disregarding the constraints and using the prior distribution in (3.1) with $(\mu_0 = 12.5, \sigma_0^2 = 3.8^2, \kappa_0 = 1, \nu_0 = 1)$. For this release, S^{2*} is larger than the nonprivate

²Note that μ_0 is more than three prior standard deviations away from the boundary of zero. Thus, one can interpret μ_0 as approximately representing a point estimate for the analyst’s prior beliefs about μ . If the analyst’s prior mean happens to be close to a boundary, the analyst likely would want to use a different prior distribution than the Gaussian. Similarly, as σ_0^2 is not close to the boundary, we can interpret it as representing approximately a point estimate for the analyst’s prior beliefs about the variance.

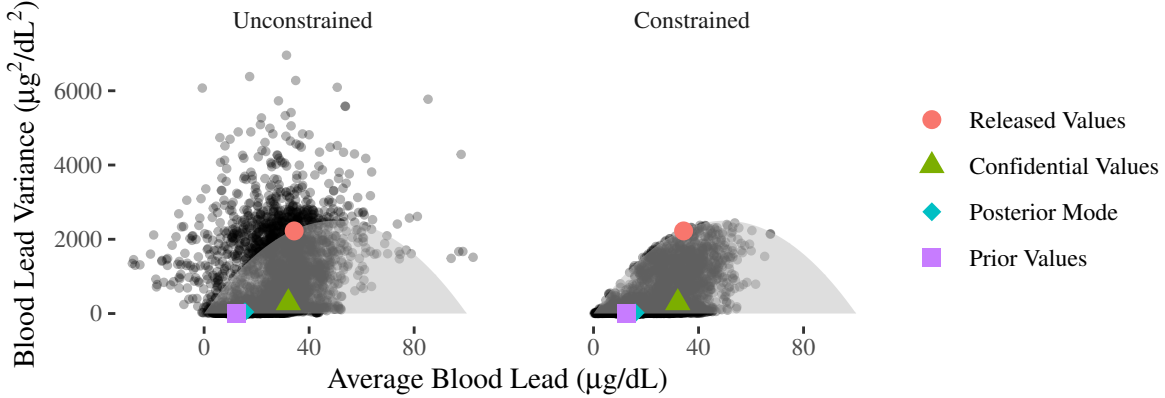


Figure 1: Joint posterior distribution for (μ, σ^2) in Example 3.6. Here, $(\bar{Y}^* = 34.30, S^{2*} = 47.17^2)$ is represented by the red circle, the unreleased $(\bar{Y} = 32.08, S^2 = 16.96^2)$ is represented by the green triangle, the analyst’s posterior mode is represented by the blue diamond, and $(\mu_0 = 12.5, S^2 = 3.8^2)$ is represented by the purple square. The left and right panels display the posterior when constraints are and are not accounted for, respectively. The shaded area represents A , the feasible region for (μ, σ^2) . Plots based on 5,000 Gibbs sampler iterations.

S^2 and is near the upper bound on σ^2 determined by Observation 3.3. Thus, more than 10% of samples from the unconstrained analysis have infeasibly large σ^2 . Additionally, the unconstrained analysis produces samples with either $\mu < 0$ or $\mu > 100$. The constrained analysis respects the bounds and produces only reasonable samples for μ and σ^2 .

In both analyses, the analyst’s posterior mode is around $\hat{\mu} \approx 16 \mu\text{g/dL}$ and $\hat{\sigma}^2 \approx 5^2 \mu\text{g}^2/\text{dL}^2$. The analyses, however, have substantially different amounts of uncertainty. The 95% HPD interval for σ^2 is $[1.2^2, 50.8^2]$ without constraints and $[1.0^2, 39.3^2]$ with constraints. Accounting for constraints allows the analyst to rule out the region with infeasibly high σ^2 , yielding a tighter interval estimate. The 95% HPD interval for μ is $[3.4, 48.2]$ without constraints and $[1.9, 42.0]$ with constraints. As above, the tighter interval in the constrained analysis is due primarily to ruling out μ in the region where σ^2 is infeasibly large. In both analyses, the posterior concentration is highest near the prior values, μ_0 and σ_0^2 ; there is much less density near the released noisy values. This indicates that the posterior inference is likely quite sensitive to the analyst’s prior distribution.

The constrained posterior is not merely a truncated version of the unconstrained posterior; truncation in this way yields inaccurate point and interval estimates. By instead truncating within the Gibbs sampler, we have coherent estimates. See Appendix E for a demonstration of the difference between the unconstrained-with-truncation posterior versus the constrained posterior.

4. DEFAULT PRIOR DISTRIBUTIONS

As evident in Example 3.6, the results of a Bayesian analysis under DP can be sensitive to the prior distribution. This is unsurprising, since the addition of noise yields data with less information about model parameters than in the public setting, making it difficult to overwhelm the information in the prior distribution. One potential fix—used in examples in

the Bayesian DP inference literature—is to increase the prior variance, yielding a weakly informative prior. One also could examine the limit in which the prior variance is infinite. We show, however, that this strategy may lead to improper limiting posterior distributions, even in settings where the nonprivate posterior is proper. In such cases, weakly informative priors may yield unreliable inference and should be avoided (Gelman, 2006).

4.1. Ensuring a Proper Posterior Distribution for the Univariate Gaussian Model Under DP. For the unconstrained analysis in the univariate Gaussian setting, a weakly informative prior of the form in (3.1) is created by making κ_0 and ν_0 small. In the limit as these hyperparameters go to zero, this produces the default prior $p(\mu, \sigma^2) \propto (\sigma^2)^{-1}$, termed the independent Jeffrey’s prior in the nonprivate setting (Sun and Berger, 2007). While it does not correspond to a proper probability distribution, the posterior distribution it produces is not only proper, but also frequentist matching. In the private setting, unfortunately, $p(\mu, \sigma^2) \propto (\sigma^2)^{-1}$ does not result in a proper posterior distribution. The following result, proved in Appendix F, demonstrates this.

Theorem 4.1. *For confidential data $Y_1, \dots, Y_n \stackrel{iid}{\sim} \mathcal{N}(\mu, \sigma^2)$ where $\bar{Y}^* \sim \text{Lap}(\bar{Y}, 1/(\varepsilon_1 n))$ and $S^{2*} \sim \text{Lap}(S^2, 1/(\varepsilon_2 n))$ are released, if an analyst has prior $p(\mu, \sigma^2) \propto (\sigma^2)^{-1}$, then $p(\mu, \sigma^2 \mid \bar{Y}^*, S^{2*})$ is not a proper probability distribution.*

An alternative strategy for finding a default prior for this setting is to focus on the likelihood’s Laplace portion. We show in Appendix G that for a Laplace likelihood in the nonprivate setting, a uniform prior produces a proper, frequentist matching posterior distribution. A uniform prior on (μ, σ^2) thus may be a reasonable choice in the private setting. Theorem 4.2, proved in Appendix F, demonstrates that $p(\mu, \sigma^2) \propto 1$ does indeed produce a proper posterior distribution.

Theorem 4.2. *For confidential data $Y_1, \dots, Y_n \stackrel{iid}{\sim} \mathcal{N}(\mu, \sigma^2)$ where $n > 3$ and $\bar{Y}^* \sim \text{Lap}(\bar{Y}, 1/(\varepsilon_1 n))$ and $S^{2*} \sim \text{Lap}(S^2, 1/(\varepsilon_2 n))$ are released, if an analyst has prior $p(\mu, \sigma^2) \propto 1$, then $p(\mu, \sigma^2 \mid \bar{Y}^*, S^{2*})$ is a proper probability distribution.*

Similarly, we can determine posterior propriety for the constrained analysis. Since a constrained uniform distribution has bounded support, it is a proper prior distribution and thus produces a proper posterior distribution. Likewise, it can be verified empirically that using the prior $p(\mu, \sigma^2) \propto (\sigma^2)^{-1}$ does not result in a proper posterior distribution in the constrained analysis.³

To assess the implications of using different prior distributions, we consider repeated sampling properties of the posterior inferences, including RMSE of point estimators and coverage rates of the posterior intervals. We do not expect the Bayesian intervals to achieve nominal coverage rates over all scenarios—none of the ones we examine do—however, these metrics are useful for measuring the impact of different prior specifications.

In particular, for each of 10,000 simulated datasets of a given sample size n , we create 95% HPD intervals based on priors that use or disregard the constraints on the data values. The Gibbs samplers are described in Appendix C.1, and run-time details are in Appendix J. Figure 2 summarizes the results. Without enforcing constraints, in this simulation scenario, the 95% HPD interval has approximately the nominal 95% coverage rate for all n considered,

³We do not prove this theoretically here, though the intuition is similar to that of Theorem 4.1.

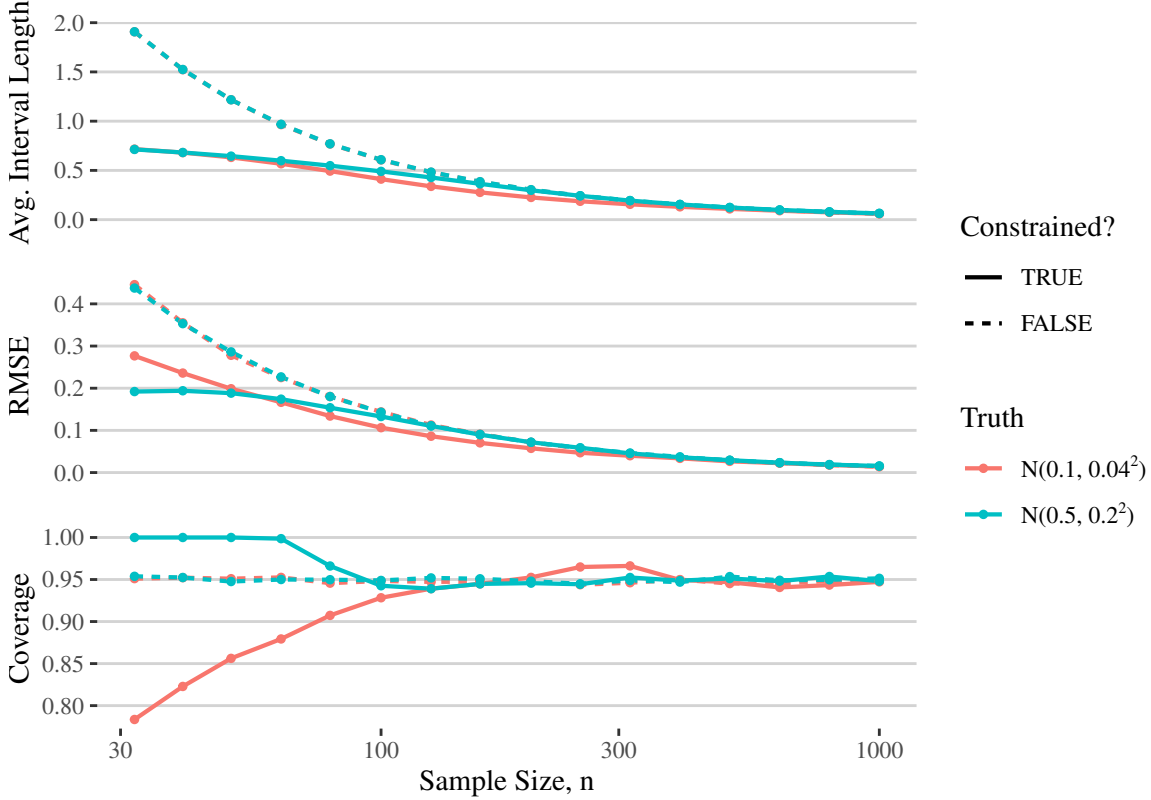


Figure 2: Average 95% HPD interval length for μ (top), root mean square error for estimating μ (middle), and empirical coverage rate of 95% HPD intervals for μ (bottom) for different sample sizes. Analyses that account for constraints (solid lines) use the uniform prior distribution over the feasible region A . Analyses that do not account for constraints (dashed lines) use the improper prior $p(\mu, \sigma^2) \propto 1$. Note that the red dashed line is below the blue dashed line. The data generating model is either $\mathcal{N}(\mu = 0.1, \sigma^2 = 0.04^2)$ (red) or $\mathcal{N}(\mu = 0.5, \sigma^2 = 0.2^2)$ (blue). DP implemented by Laplace mechanisms with $\varepsilon_1 = \varepsilon_2 = 0.1$ and bounding $Y_i \in [0, 1]$. Each Gibbs sampler is run for 20,000 iterations. Results based on 10,000 simulated datasets.

and both the average interval length and RMSE decay proportionally to $1/n$.⁴ When enforcing the constraints, the results are similar to those for the unconstrained analysis when n is large. In these cases, credible intervals for μ are far from the boundary, so the constraints have little effect. Additionally, when n is large, the prior distributions exert less influence on the posterior distributions. When n is small, the 95% HPD interval has lower than 95% coverage rate when the true μ is close to the boundary and near 100% coverage when μ is in the middle of the range. Since the DP statistics have modest information content when n is small (especially when ϵ is small), the prior distribution can have a substantial

⁴A regression of log average interval length on log n has a slope of -1 (with $R^2 > 0.999$), indicating that average interval length decays proportionally to $1/n$. A similar result holds for the RMSE.

influence on the Bayesian inference, in that the posterior distribution is pulled towards the prior distribution. In fact, it can be shown that the uniform prior over the region where $\mu \in [0, 1]$ and $\sigma^2 \in (0, \mu(1 - \mu)]$ induces the marginal distribution $\mu \sim \text{Beta}(2, 2)$. This Beta distribution places more probability density in the center of the distribution than near 0 or 1. Thus, when $\mu = 0.5$, the prior distribution puts high mass in regions around the true population parameter(s), resulting in posterior intervals that have higher than 95% coverage rates. When $\mu = 0.1$, the coverage rates are below 95%. This suggests that the proper uniform prior distribution is most appropriate when an analyst believes μ more likely is near the center of the distribution than the tails.

Enforcing the constraint has advantages. For all n in Figure 2, the average interval length for the constrained analysis is less than 0.75. Meanwhile, the average interval length without constraints is greater than 1 for $n \leq 50$, indicating that the credible intervals must include values that are not within $[0, 1]$. Such intervals often include all of $[0, 1]$; these are practically useless, as they tell the analyst nothing about the location of the parameter values. In these scenarios, the unconstrained analysis achieves calibrated coverage rates at the high cost of inactionable inference. Additionally, the middle panel of Figure 2 presents the RMSE from estimating μ with the posterior mode. We see that the RMSE is uniformly lower in the constrained analysis than in the unconstrained analysis. Summarizing, the incorporation of constraints yields estimates that are more sensible scientifically, have less error, and have tighter interval estimates. However, these advantages come with a tradeoff: because of the influence of the prior distribution, the resulting intervals may have less than nominal frequentist coverage rates, especially when n is small and μ is near the boundary.

Figure 3 presents the results of a simulation study identical to that used in Figure 2, but with $\varepsilon = 2$. There is little practical difference between the constrained and unconstrained analysis at this larger ε . The unconstrained analysis no longer offers approximately nominal coverage rates: the 95% posterior intervals cover more than 95% of the time. To explain this, we first note that the Laplace noise tends to dwarf the Gaussian noise when $\varepsilon = 0.2$. Since $p(\mu, \sigma^2) \propto 1$ is frequentist-matching to a Laplace likelihood, it results in a coverage rate that matches the nominal 95% rate (albeit with sometimes nonsensical intervals). In contrast, when $\varepsilon = 2$, both the Gaussian and Laplace kernels contribute meaningfully to the likelihood function, and the prior distribution no longer enjoys the frequentist-matching property.

4.2. Default Priors for Example 3.6. Returning to the setting of Example 3.6, we examine the effect of replacing the informative prior used in Section 3.2.2 with the uniform default priors. Figure 4 presents a plot analogous to Figure 1 but now using $p(\mu, \sigma^2) \propto 1$ for the unconstrained analysis and the uniform prior over A , as defined in (3.4), for the constrained analysis. In the unconstrained analysis, the posterior mode is approximately equal to the released statistics, but more than 50% of posterior draws are outside of the feasible region. In the constrained analysis, as discussed above, the posterior mode for μ is shifted towards the center of the distribution. A similar effect is observed for σ^2 ; the prior has more mass closer to zero and so the posterior mode is smaller than S^{2*} .

Analysts also might be interested in using the Bayesian model for prediction. Figure 5 plots the posterior predictive distributions for the two default priors. To simulate these distributions, at each Gibbs iteration t we sample a new $Y_{\text{new}(t)} \sim \mathcal{N}(\mu_{(t)}, \sigma_{(t)}^2)$ for the unconstrained analysis and a new $Y_{\text{new}(t)} \sim \mathcal{N}_{[0,1]}(\mu_{(t)}, \sigma_{(t)}^2)$ for the constrained analysis. The

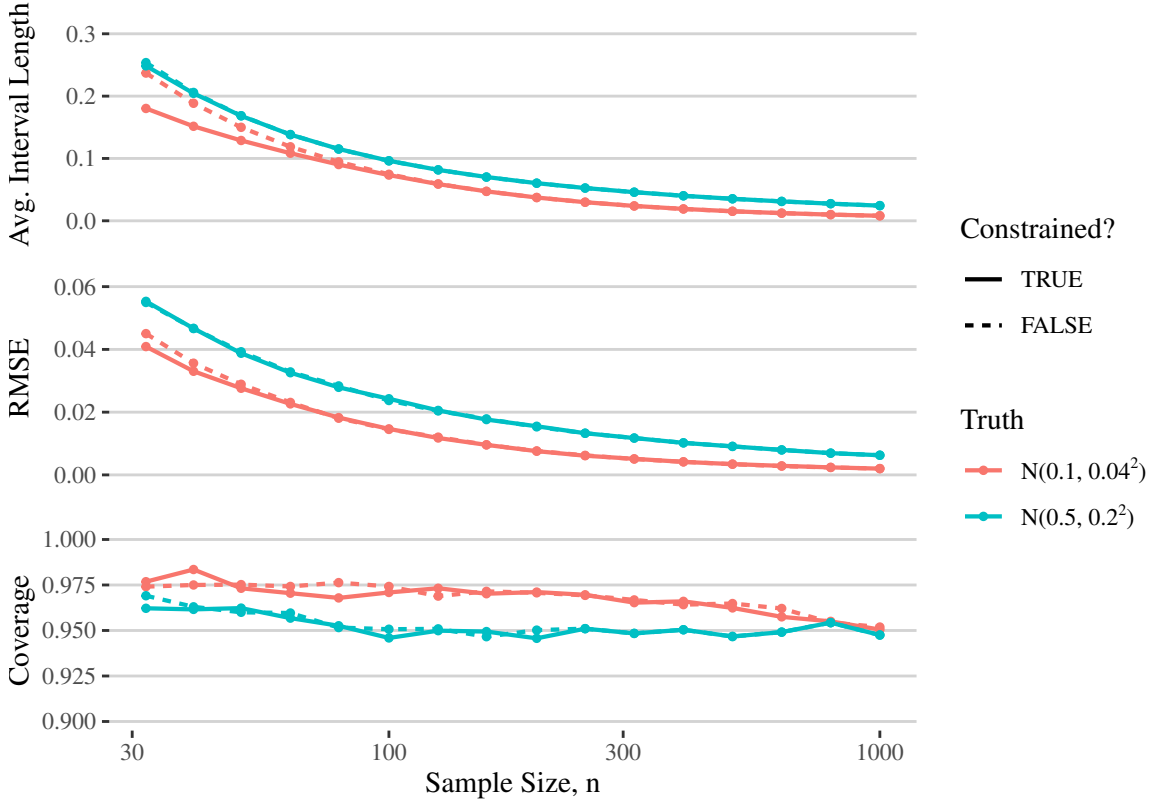


Figure 3: Repeating the simulation from Figure 2 but with $\varepsilon_1 = \varepsilon_2 = 1$.

constrained posterior predictive distribution only includes values in the feasible region of 0 to 100 $\mu\text{g/dL}$, with the distribution centered around the posterior mode. The unconstrained posterior predictive distribution, meanwhile, generates 24% of predictive draws as negative values and 10% as values greater than 100 $\mu\text{g/dL}$. This leads to the constrained analysis having substantially lower predictive variance than the unconstrained: the standard deviation is 53 $\mu\text{g/dL}$ without constraints and 25 $\mu\text{g/dL}$ with constraints.

To obtain reasonable predictions with the draws from the unconstrained analysis, an analyst must resort to ad hoc methods. They might, for example, re-code all negative predictions to 0 and all predictions greater than 100 to 100. This would lead to 24% of predictions being exactly zero, which is not reasonable scientifically, and results in an inflated predictive standard deviation of 35 $\mu\text{g/dL}$. Alternatively, the analyst could sample predictions from a truncated Gaussian distribution. The variance of this distribution is still inflated by the samples of overly large values of $\sigma_{(t)}^2$, leading to a predictive standard deviation of 27 $\mu\text{g/dL}$ under this ad hoc approach, which is larger than that of the theoretically principled constrained analysis.

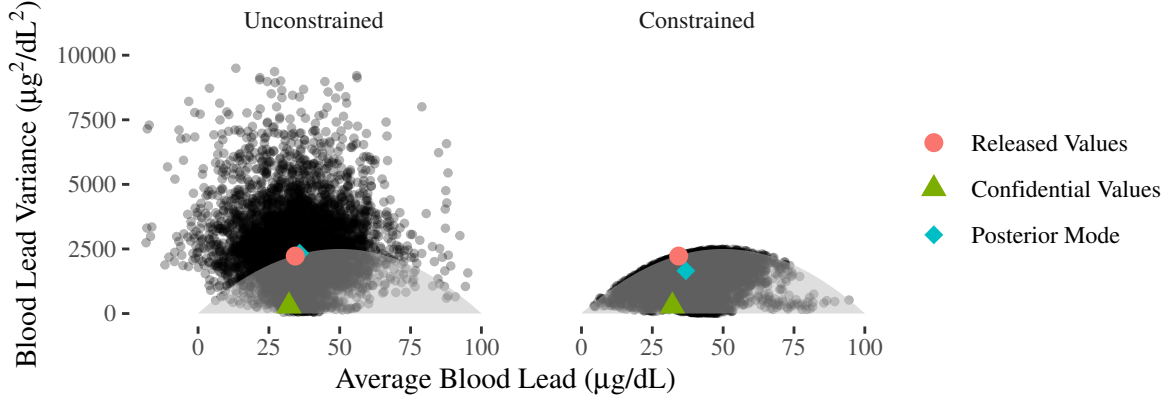


Figure 4: Posterior distribution for (μ, σ^2) in Example 3.6 under default priors. The point $(\bar{Y}^* = 34.30, S^2 = 47.17^2)$ is represented by the red circle; the unreleased point $(\bar{Y} = 32.08, S^2 = 16.96^2)$ is represented by the green triangle; and, the analyst's posterior mode is represented by the blue diamond. The left and right panels provide the posterior when constraints are and are not accounted for, respectively. The shaded area represents A , the feasible region for (μ, σ^2) . This plot is based on 5,000 Gibbs iterations.

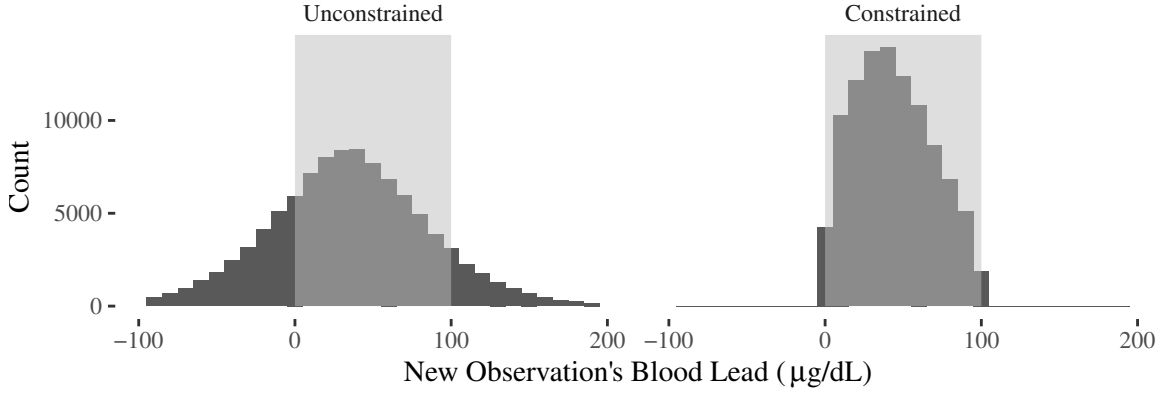


Figure 5: Posterior predictive distribution of a new observation for the posterior draws from Example 3.6 using the default prior distributions. The left and right panels provide the posterior predictive distributions without and with accounting for constraints, respectively. The shaded area represents the feasible region for a new observation. Plots based on 100,000 Gibbs sampler iterations.

5. REGRESSION APPLICATION

The strategies to account for constraints on data values can be generalized to settings beyond the univariate Gaussian. We demonstrate such an extension using the DP Bayesian linear regression method of [Bernstein and Sheldon \(2019\)](#) and their worked example.

Example 5.1. Researchers wish to estimate the association of a state’s per capita drinking rate, $\{x_i\}$, with the state’s cirrhosis death rate, $\{y_i\}$, using data from [Brownlee \(1965\)](#). The data comprise $n = 46$ observations. The ranges of both variables are public information.

In [Bernstein and Sheldon \(2019\)](#), the authors rescale both variables to lie in $[0, 1]$. Letting $\mathbf{Y}$ be the vector of responses and $\mathbf{X}$ be a design matrix with a column of ones, the authors’ model is a linear regression of the form $\mathbf{Y} \sim \mathcal{N}_n(\mathbf{X}\boldsymbol{\theta}, \sigma^2 \mathbf{I}_n)$ with the conjugate prior

$$\boldsymbol{\theta} \mid \sigma^2 \sim \mathcal{N}_2 \left(\begin{bmatrix} 1 \\ 0 \end{bmatrix}, \frac{\sigma^2}{4} \mathbf{I}_2 \right), \quad \sigma^2 \sim \text{IG} \left(\frac{40}{2}, \frac{1}{2} \right), \quad (5.1)$$

where $\mathbf{I}_p$ is the $p \times p$ identity matrix. Their proposed method applies a Laplace mechanism to each element of the sufficient statistics $\mathbf{X}^\top \mathbf{X}$, $\mathbf{X}^\top \mathbf{Y}$ and $\mathbf{Y}^\top \mathbf{Y}$. They force the matrix $\mathbf{Z}^\top \mathbf{Z}$, where $\mathbf{Z} = [\mathbf{X} \ \mathbf{Y}]$, to be positive semi-definite (PSD) by projecting samples to the nearest PSD matrix. They estimate the posterior distribution of the regression parameters via their Gibbs-SS-Noisy method. This method also adds Laplace noise to the matrix of fourth sample moments of $\mathbf{X}$.

We utilize methods identical to those in [Bernstein and Sheldon \(2019\)](#) with the following modifications. Since it is known that $x_i, y_i \in [0, 1]$ for all i , we can exploit the structure of the linear regression to constrain parameter updates in the Gibbs sampler. In particular, letting $\mathbf{X} = [\mathbf{1}_n \ \mathbf{x}_1]$, we show in [Appendix H](#) that

$$0 \leq \mathbf{x}_1^\top \mathbf{x}_1 \leq \mathbf{x}_1^\top \mathbf{1} \leq n \quad (5.2)$$

$$0 \leq \mathbf{Y}^\top \mathbf{Y} \leq \mathbf{Y}^\top \mathbf{1} \leq n \quad (5.3)$$

$$\mathbf{x}_1^\top \mathbf{Y} \leq \min\{\mathbf{x}_1^\top \mathbf{1}, \mathbf{Y}^\top \mathbf{1}\} \quad (5.4)$$

$$0 \leq \sigma^2 \leq 1/4. \quad (5.5)$$

Letting the regression coefficients be $\boldsymbol{\theta} = [\theta_0 \ \theta_1]^\top$, we know that (θ_1, θ_0) must lie in the region depicted in the right panels of [Figure 6](#). We create a constrained Gibbs sampler by drawing each iterate from the unconstrained full conditional distribution and rejecting and resampling whenever all constraints are not satisfied. We also reject and resample if $\mathbf{Z}^\top \mathbf{Z}$ is not PSD.

For demonstration, we draw a set of DP sufficient statistics with $\varepsilon = 0.1$ for each of the 11 queries.⁵ [Figure 6](#) displays posterior draws for two of the sufficient statistics and for the parameters, using an off-the-shelf application of the method of [Bernstein and Sheldon \(2019\)](#) and our modification that accounts for the constraints. For both samplers, the posterior concentrates around the confidential values of the sufficient statistics. However, by incorporating the constraint that $0 \leq \mathbf{Y}^\top \mathbf{Y} \leq \mathbf{Y}^\top \mathbf{1} \leq n$, the modified sampler yields a substantial reduction in the posterior variance. As indicated by the red points, nearly 20% of posterior draws without accounting for constraints do not satisfy this inequality. For the off-the-shelf application of the method of [Bernstein and Sheldon \(2019\)](#), a small percentage (less than 5%) of the posterior draws of (θ_0, θ_1) do not satisfy the constraint

⁵The addition of Laplace noise to the fourth sample moments of $\mathbf{X}$ adds additional uncertainty which is not accounted for in the procedure in [Bernstein and Sheldon \(2019\)](#). Because of this, our rejection sampling approach struggles when the released noisy moments are far from their nonprivate counterparts, leading to the Gibbs sampling procedure being unable to find a draw for the next Gibbs iterate satisfying the constraints and getting “stuck.” One could avoid this issue by truncating the full conditionals using methods akin to those described in [Section 3.2](#), or by simply running the sampler for long enough that it is able to produce a proposal that satisfies the constraints. For ease of reproducibility, [Figure 6](#) utilizes a draw for the sample moments of $\mathbf{X}$ where this issue is not encountered.

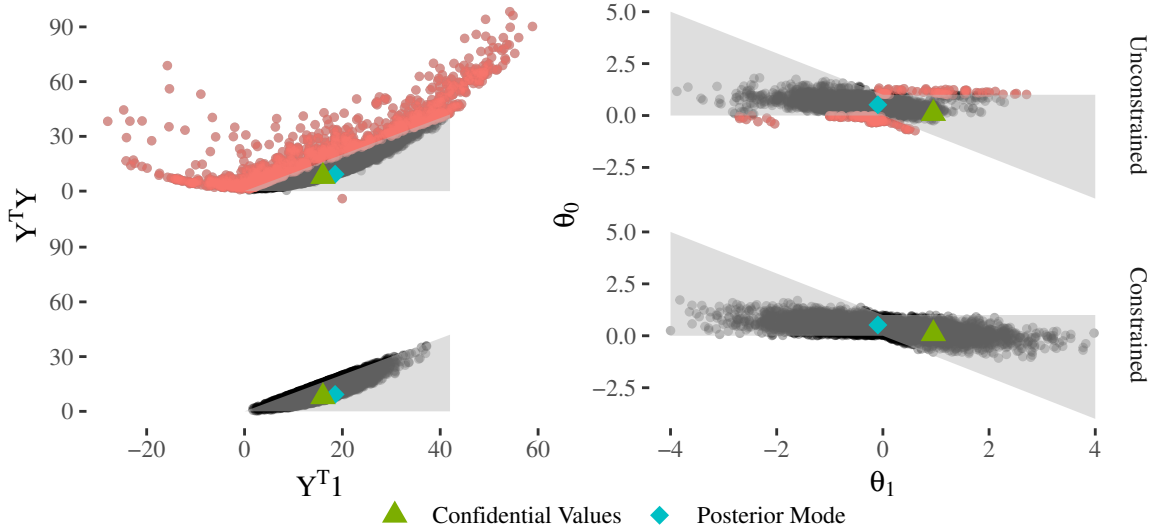


Figure 6: Posterior draws using the method of [Bernstein and Sheldon \(2019\)](#) for estimating DP linear regressions. The left panels represent draws of the imputed sufficient statistics $\mathbf{Y}^\top \mathbf{1}$ and $\mathbf{Y}^\top \mathbf{Y}$, while the right panels represent draws of the parameters θ_1 and θ_0 . The upper and lower panels provide the posterior when constraints are and are not accounted for, respectively. The shaded areas represent the feasible regions; points outside the feasible regions are colored in red. In the left panels, the green triangle represents the confidential values of $\mathbf{Y}^\top \mathbf{1}$ and $\mathbf{Y}^\top \mathbf{Y}$. In the right panels, the green triangle represents the nonprivate OLS estimator of (θ_0, θ_1) . In all panels, the blue diamond represents the posterior mode. This plot is based on 10,000 Gibbs iterations.

for $\boldsymbol{\theta}$. Notably, when all relevant constraints are enforced, the posterior distribution is more spread out throughout the feasible region. Without constraints, the nonprivate OLS estimate for θ_1 —represented by the green triangle—is not within a 95% interval estimate for the parameter. The increased variance due to the constraints, however, fixes this issue. Fitting this model without constraints leads artificially to a model with too little uncertainty about its parameter estimates, which subsequently can lead to underestimated uncertainty in interval estimates and downstream predictions.

6. DISCUSSION

The Bayesian paradigm is appealing for post-processing inference with DP, in that it facilitates convenient propagation of uncertainty. However, as in all Bayesian inference settings, the posterior distribution can be sensitive to the analyst’s prior specification. Our work shows that this sensitivity extends to the analyst’s decision whether or not the prior distribution should account for the constraints on the parameters implied by enforcing bounds on the data values. Our simulation studies suggest that incorporating constraints in Bayesian DP inference can result in posterior distributions with support only on feasible parameter values, as well as more accurate point estimates and predictions. Our work also suggests default prior distributions that ensure proper posterior distributions. Taken

together, the results suggest scope for further development of default prior distributions that lead to (proper) posterior distributions with desirable frequentist properties.

We modify the Gibbs sampler of [Bernstein and Sheldon \(2018\)](#) primarily as a step along the way to obtaining computationally efficient Bayesian inference under constraints. Nonetheless, we anticipate that sampling from exact full conditional distributions can have benefits in practical applications involving Gaussian models. In addition to the examples mentioned previously, Gaussian distributions can be used in a variety of tasks in statistical analysis, including A/B testing and hierarchical modeling.

Although the bounds for μ and σ^2 are guaranteed to hold without any need to peek at the confidential data, they are likely to be conservative for data that follow Gaussian distributions. For example, with the values of n used in our simulation scenarios, it is unlikely that values of the confidential data will be more than three or four standard deviations from μ . The analyst could define bounds that accord with this feature of the Gaussian distribution. To illustrate, suppose that $\mu = 0.1$ and $n = 100$. For the Gaussian model to be plausible, we would expect σ^2 to be at most $0.1/4 = .025$, since we do not expect observations to be more than four standard deviations away from μ . This is a stronger bound than the .09 resulting from Observation 3.3. This example suggests that analysts might be able to set prior distributions with tighter bounds while still maintaining the computational simplicity and convenience of the measurement error model framework for Bayesian inference. There are likely other strategies that could be employed to leverage the normality assumption; investigating specifications for such prior distributions is an interesting topic for future work.

There are settings where the underlying data values are bounded but a Gaussian assumption is not appropriate. In such settings, presuming a Gaussian model could result in unreliable posterior inferences. Nonetheless, we expect the general approach of enforcing constraints on the sufficient statistics and parameter values via measurement error modeling can be beneficial for Bayesian inference for other models. When convenient full conditionals are not available for Gibbs samplers, it may be possible to adapt the general sampling strategies proposed by [Bernstein and Sheldon \(2018\)](#) or [Ju et al. \(2022\)](#)—see Appendix I for an illustrative adaptation of the latter for the univariate Gaussian model—and possibly other MCMC sampling strategies, to incorporate constraints. For example, analysts can follow the strategy used for linear regression in Section 5, i.e., (1) determine implied constraints on the statistics and/or parameters and (2) incorporate the constraints into the sampler by truncating the appropriate full conditionals or rejecting draws outside of the constrained region.

To illustrate how the general strategy can apply in other models, we present the following two settings.

- (1) Suppose an analyst is interested in inferring the rate of incoming phone calls for a company. Let T_1 represent the time from the start of the day to the first phone call, T_2 represent the time between the first and second call, T_3 represent the time between the second and third call, etc. A common assumption is that $T_1, \dots, T_n \sim \text{Exp}(\lambda)$, i.e., an exponential distribution where λ is the parameter of interest. Suppose calls are typically short and frequent, and each wait time can be reasonably bounded above by, say, 60 minutes. The analyst is provided a noisy version $\bar{T}^*$ of the sample mean $\bar{T}$ via the Laplace mechanism, computed using the bound of 60 minutes. It is straightforward to show that the bounds on the data imply that $0 \leq \bar{T} \leq 60$ and $\lambda \geq 1/60$. These constraints can be incorporated into a Gibbs sampler via truncation of the prior distribution and the appropriate full conditionals.

- (2) Suppose an analyst is interested in estimating the rate of mutations within a particular part of a DNA strand. Let N_i denote the number of observed mutations for subject i and θ denote the rate of mutation. A common assumption is that $N_1, \dots, N_n \sim \text{Pois}(\theta)$, i.e., a Poisson distribution where θ is the rate parameter. An expert recommends that the number of mutations is reasonably bounded above by 20. The analyst has a noisy version $\bar{N}^*$ of the sample mean $\bar{N}$ released via the geometric mechanism. The bounds imply that $0 \leq \bar{N} \leq 20$ and $0 \leq \theta \leq 20$. These constraints can be incorporated in the prior distribution and in an accept/reject step for the appropriate full conditionals in a Gibbs sampler.

While these examples involve single-parameter models, the strategies for incorporating bounds on data values in Bayesian inference also may be applicable in more complex models. For example, if an explanatory variable in logistic regression is known to have range $[0, 1]$, constraints on the statistics analogous to those derived for linear regression in Appendix H will apply and can be incorporated into a Gibbs sampling scheme. Further analysis of the generalization of these ideas from the Gaussian setting to non-Gaussian and multivariate settings is a promising avenue for future work.

ACKNOWLEDGMENT

This research was supported by NSF grant SES-2217456.

REFERENCES

- G. Bernstein and D. R. Sheldon. Differentially private Bayesian inference for exponential families. *Advances in Neural Information Processing Systems*, 31, 2018. <https://dl.acm.org/doi/10.5555/3327144.3327215>.
- G. Bernstein and D. R. Sheldon. Differentially private Bayesian linear regression. *Advances in Neural Information Processing Systems*, 32, 2019. <https://dl.acm.org/doi/10.5555/3454287.3454335>.
- S. Biswas, Y. Dong, G. Kamath, and J. Ullman. Coinpress: Practical private mean and covariance estimation. *Advances in Neural Information Processing Systems*, 33, 2020. <https://dl.acm.org/doi/10.5555/3495724.3496937>.
- K. A. Brownlee. Statistical theory and methodology in science and engineering. <https://people.sc.fsu.edu/~jburkardt/datasets/regression/x20.txt>, 1965. Accessed on 27 January 2025.
- S. Couch, Z. Kazan, K. Shi, A. Bray, and A. Groce. Differentially private nonparametric hypothesis testing. In *Proceedings of the 2019 ACM SIGSAC Conference on Computer and Communications Security*, pages 737–751. Association for Computing Machinery, 2019. <https://doi.org/10.1145/3319535.3339821>.
- D. M. Diez, C. D. Barr, and M. Cetinkaya-Rundel. *OpenIntro Statistics*, volume 4. OpenIntro Boston, MA, USA:, 2012.
- C. Dimitrakakis, B. Nelson, Z. Zhang, A. Mitrokotsa, and B. I. Rubinstein. Differential privacy for Bayesian inference through posterior sampling. *Journal of Machine Learning Research*, 18(11):1–39, 2017. <https://dl.acm.org/doi/abs/10.5555/3122009.3122020>.
- V. D’Orazio, J. Honaker, and G. King. Differential privacy for social science inference, 2015. Sloan Foundation Economics Research Paper 2676160.

- W. Du, C. Foot, M. Moniot, A. Bray, and A. Groce. Differentially private confidence intervals. *arXiv preprint arXiv:2001.02285*, 2020. <https://doi.org/10.48550/arXiv.2001.02285>.
- C. Dwork, F. McSherry, K. Nissim, and A. Smith. Calibrating noise to sensitivity in private data analysis. In *Theory of Cryptography: Third Theory of Cryptography Conference, TCC 2006*, volume 3876. Springer, 2006. https://doi.org/10.1007/11681878_14.
- C. Ferrando, S. Wang, and D. Sheldon. Parametric bootstrap for differentially private confidence intervals. In *Proceedings of The 25th International Conference on Artificial Intelligence and Statistics*, volume 151, pages 1598–1618. PMLR, 2022. <https://proceedings.mlr.press/v151/ferrando22a.html>.
- A. Gelman. Prior distributions for variance parameters in hierarchical models (comment on article by Browne and Draper). *Bayesian Analysis*, 1(3):515–534, 2006. <https://doi.org/10.1214/06-BA117A>.
- J. Geumlek, S. Song, and K. Chaudhuri. Rényi differential privacy mechanisms for posterior sampling. *Advances in Neural Information Processing Systems*, 30, 2017. https://proceedings.neurips.cc/paper_files/paper/2017/file/56584778d5a8ab88d6393cc4cd11e090-Paper.pdf.
- R. Gong. Exact inference with approximate computation for differentially private data via perturbations. *Journal of Privacy and Confidentiality*, 12(2), 2022. <https://doi.org/10.29012/jpc.797>.
- P. D. Hoff. *A First Course in Bayesian Statistical Methods*, volume 580. Springer, 2009.
- J. E. Jewson, S. Ghalebikesabi, and C. C. Holmes. Differentially private statistical inference through beta-divergence one posterior sampling. *Advances in Neural Information Processing Systems*, 36, 2024. <https://dl.acm.org/doi/10.5555/3666122.3669487>.
- N. L. Johnson, S. Kotz, and N. Balakrishnan. *Continuous Univariate Distributions. Vol. 1 (2nd ed.)*. New York: Wiley, 1994.
- N. Ju, J. Awan, R. Gong, and V. Rao. Data augmentation MCMC for Bayesian inference from privatized data. *Advances in Neural Information Processing Systems*, 35, 2022. <https://dl.acm.org/doi/10.5555/3600270.3601195>.
- G. Kamath, A. Mouzakis, M. Regehr, V. Singhal, T. Steinke, and J. Ullman. A bias-accuracy-privacy trilemma for statistical estimation. *Journal of the American Statistical Association*, 120:2338–2349, 2025. <https://doi.org/10.1080/01621459.2024.2443275>.
- V. Karwa and S. Vadhan. Finite sample differentially private confidence intervals. In *9th Innovations in Theoretical Computer Science Conference (ITCS 2018). Leibniz International Proceedings in Informatics (LIPIcs)*, volume 94, pages 44:1–44:9. Schloss Dagstuhl – Leibniz-Zentrum für Informatik, 2018. <https://doi.org/10.4230/LIPIcs.ITCS.2018.44>.
- S. P. Kasiviswanathan and A. Smith. On the ‘semantics’ of differential privacy: A Bayesian formulation. *Journal of Privacy and Confidentiality*, 6(1), 2014. <https://doi.org/10.29012/jpc.v6i1.634>.
- Z. Kazan and J. P. Reiter. Prior-itzing privacy: A Bayesian approach to setting the privacy budget in differential privacy. *Advances in Neural Information Processing Systems*, 2024. <https://dl.acm.org/doi/10.5555/3737916.3740785>.
- Z. Kazan and J. P. Reiter. Assessing statistical disclosure risk for differentially private, hierarchical count data, with application to the 2020 US decennial census. *Statistica Sinica*, 35:629–649, 2025. <https://doi.org/10.5705/ss.202022.0187>.

- Z. Kazan, K. Shi, A. Groce, and A. P. Bray. The test of tests: A framework for differentially private hypothesis testing. In *Proceedings of the 40th International Conference on Machine Learning*, pages 16131–16151. PMLR, 2023. <https://dl.acm.org/doi/10.5555/3618408.3619068>.
- D. Kifer and A. Machanavajjhala. Pufferfish: A framework for mathematical privacy definitions. *ACM Transactions on Database Systems (TODS)*, 39(1), 2014. <https://doi.org/10.1145/2514689>.
- D. Kifer, J. M. Abowd, R. Ashmead, R. Cumings-Menon, P. Leclerc, A. Machanavajjhala, W. Sexton, and P. Zhuravlev. Bayesian and frequentist semantics for common variations of differential privacy: Applications to the 2020 census. Technical report, U.S. Bureau of the Census, CED-WP-2022-004, 2022. <https://doi.org/10.48550/arXiv.2209.03310>.
- A. Kulesza, A. T. Suresh, and Y. Wang. Mean estimation in the add-remove model of differential privacy. In *Proceedings of the 41st International Conference on Machine Learning*. PMLR, 2024. <https://dl.acm.org/doi/10.5555/3692070.3693097>.
- T. Kulkarni, J. Jälkö, A. Koskela, S. Kaski, and A. Honkela. Differentially private Bayesian inference for generalized linear models. In *Proceedings of the 38th International Conference on Machine Learning*, pages 5838–5849. PMLR, 2021. <https://proceedings.mlr.press/v139/kulkarni21a/kulkarni21a.pdf>.
- L. Li and J. P. Reiter. Bayesian inference for estimating subset proportions using differentially private counts. *Journal of Survey Statistics and Methodology*, 10(3):785–803, 2022. <https://doi.org/10.1093/jssam/smab060>.
- D. McClure and J. P. Reiter. Differential privacy and statistical disclosure risk measures: An investigation with binary synthetic data. *Transactions on Data Privacy*, 5(3):535–552, 2012. <https://dl.acm.org/doi/10.5555/2423656.2423658>.
- K. Minami, H. Arai, I. Sato, and H. Nakagawa. Differential privacy without sensitivity. In D. Lee, M. Sugiyama, U. Luxburg, I. Guyon, and R. Garnett, editors, *Advances in Neural Information Processing Systems*, volume 29. Curran Associates, Inc., 2016. <https://dl.acm.org/doi/10.5555/3157096.3157204>.
- W. Mortada, M. Sobh, M. M. El-Defrawy, and S. E. Farahat. Study of lead exposure from automobile exhaust as a risk for nephrotoxicity among traffic policemen. *American Journal of Nephrology*, 21(4):274–279, 2001. <https://doi.org/10.1159/000046261>.
- F. Novomestky and S. Nadarajah. *truncdist: Truncated Random Variables*, 2016. <https://CRAN.R-project.org/package=truncdist>. R package version 1.0-2.
- T. Park and G. Casella. The Bayesian lasso. *Journal of the American Statistical Association*, 103(482):681–686, 2008. <https://doi.org/10.1198/016214508000000337>.
- V. Peña and A. F. Barrientos. Differentially private hypothesis testing with the subsampled and aggregated randomized response mechanism. *Statistica Sinica*, 35:1–21, 2025. <https://doi.org/10.5705/ss.202022.0279>.
- O. Sheffet. Differentially private ordinary least squares. In *Proceedings of the 34th International Conference on Machine Learning*, volume 70, pages 3105–3114. PMLR, 2017. <https://dl.acm.org/doi/10.5555/3305890.3306002>.
- D. Sun and J. O. Berger. Objective Bayesian analysis for the multivariate normal model. In *Bayesian Statistics 8: Proceedings of the Eighth Valencia International Meeting June 2–6, 2006*. Oxford University Press, 2007. <https://doi.org/10.1093/oso/9780199214655.003.0020>.
- Y.-X. Wang. Revisiting differentially private linear regression: optimal and adaptive prediction & estimation in unbounded domain. In *Conference on Uncertainty in Artificial*

- Intelligence*, 2018. <https://www.auai.org/uai2018/proceedings/papers/40.pdf>.
- Y.-X. Wang, S. Fienberg, and A. Smola. Privacy for free: Posterior sampling and stochastic gradient Monte Carlo. In *Proceedings of the 32nd International Conference on Machine Learning*, volume 37, pages 2493–2502. PMLR, 2015. <https://dl.acm.org/doi/10.5555/3045118.3045383>.
- J. Zhang, Z. Zhang, X. Xiao, Y. Yang, and M. Winslett. Functional mechanism: regression analysis under differential privacy. *Proceedings of the VLDB Endowment*, 5(11):1364–1375, 2012. <https://doi.org/10.14778/2350229.2350253>.
- W. Zhang and R. Zhang. Dp-fast mh: Private, fast, and accurate Metropolis-Hastings for large-scale Bayesian inference. In *Proceedings of the 40th International Conference on Machine Learning*. PMLR, 2023. <https://dl.acm.org/doi/10.5555/3618408.3620166>.
- Z. Zhang, B. I. Rubinstein, and C. Dimitrakakis. On the differential privacy of Bayesian inference. In D. Schuurmans and M. P. Wellman, editors, *Proceedings of the Thirtieth AAAI Conference on Artificial Intelligence, February 12-17, 2016, Phoenix, Arizona, USA*, pages 2365–2371. AAAI Press, 2016. <https://doi.org/10.1609/aaai.v30i1.10254>.

APPENDIX A. RE-SCALING THE DATA

Throughout the article, we assume that the data $Y_1, \dots, Y_n$ are contained in the interval $[a, b]$ for known, public a and b . In this section, we prove that, without loss of generality, we may consider data on the interval $[0, 1]$ by re-scaling the data as follows to obtain $\tilde{Y}_i \in [0, 1]$,

$$\tilde{Y}_i = \frac{Y_i - a}{b - a}. \quad (\text{A.1})$$

We can convert back to the original scale via the relationship $Y_i = (b - a)\tilde{Y}_i + a$.

We now show that the sufficient statistics released via the Laplace mechanism on the $[0, 1]$ scale are re-scaled versions of the sufficient statistics released via the Laplace mechanism on the $[a, b]$ scale. Thus no information is lost from scaling before the DP release and re-scaling afterwards.

Observation 3.1. *Let $Y_1, \dots, Y_n \in [a, b]$ and let $\tilde{Y}_i = (Y_i - a)/(b - a) \in [0, 1]$. Let $\bar{Y}$ and S^2 be the sample mean and variance for $\{Y_i\}$ and let $\tilde{\bar{Y}}$ and $\tilde{S}^2$ be the sample mean and variance for $\{\tilde{Y}_i\}$. Suppose each statistic is released via the Laplace mechanism under ε -DP and denote the DP statistics $\bar{Y}^*$, S^{2*} , $\tilde{\bar{Y}}^*$, and $\tilde{S}^{2*}$. Then, $\bar{Y}^* \stackrel{d}{=} (b - a)\tilde{\bar{Y}}^* + a$ and $S^{2*} \stackrel{d}{=} (b - a)^2\tilde{S}^{2*}$.*

Proof. To begin, note that we may relate the sample means on the two scales as follows

$$\bar{Y} = \frac{1}{n} \sum_{i=1}^n Y_i = \frac{1}{n} \sum_{i=1}^n [(b - a)\tilde{Y}_i + a] = (b - a) \frac{1}{n} \sum_{i=1}^n \tilde{Y}_i + a = (b - a)\tilde{\bar{Y}} + a. \quad (\text{A.2})$$

Similarly, we may relate the sample variances on the two scales as follows.

$$S^2 = \frac{1}{n - 1} \sum_{i=1}^n (Y_i - \bar{Y})^2 \quad (\text{A.3})$$

$$= \frac{1}{n - 1} \sum_{i=1}^n ([(b - a)\tilde{Y}_i + a] - [(b - a)\tilde{\bar{Y}} + a])^2 \quad (\text{A.4})$$

$$= (b - a)^2 \frac{1}{n - 1} \sum_{i=1}^n (\tilde{Y}_i - \tilde{\bar{Y}})^2 \quad (\text{A.5})$$

$$= (b - a)^2 \tilde{S}^2. \quad (\text{A.6})$$

We now consider the sensitivities of each of the four statistics. By Lemmas 11 and 12 in [Du et al. \(2020\)](#), the sensitivity of $\bar{Y}$ is $(b - a)/n$ and the sensitivity of S^2 is $(b - a)^2/n$. Similarly, the sensitivity of $\tilde{\bar{Y}}$ is $1/n$ and the sensitivity of $\tilde{S}^2$ is $1/n$. Thus, when released under ε -DP via the Laplace mechanism, the released statistics have distribution

$$\bar{Y}^* \sim \text{Lap}\left(\bar{Y}, \frac{b - a}{\varepsilon n}\right), \quad S^{2*} \sim \text{Lap}\left(S^2, \frac{(b - a)^2}{\varepsilon n}\right) \quad (\text{A.7})$$

$$\tilde{\bar{Y}}^* \sim \text{Lap}\left(\tilde{\bar{Y}}, \frac{1}{\varepsilon n}\right), \quad \tilde{S}^{2*} \sim \text{Lap}\left(\tilde{S}^2, \frac{1}{\varepsilon n}\right). \quad (\text{A.8})$$

Recall that if X has a Laplace distribution with location parameter m and scale parameter b , i.e., $X \sim \text{Lap}(m, b)$, then if $k > 0$ and $c \in \mathbb{R}$, it follows that $kX + c \sim \text{Lap}(km + c, kb)$. By this property, $(b - a)\tilde{\bar{Y}}^* + a$ has location parameter $(b - a)\tilde{\bar{Y}} + a = \bar{Y}$ and scale parameter

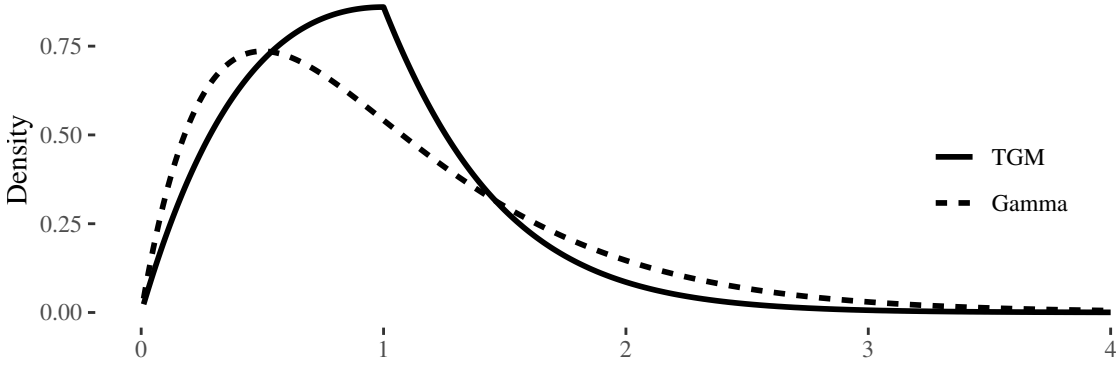


Figure 7: Comparison of the probability density functions of $\mathbf{TGM}(\alpha = 2, \beta = 2, \lambda = 1, \tau = 1)$ (the solid line) and $\mathbf{Gam}(\alpha = 2, \beta = 2)$ (the dashed line).

$(b - a)/(\varepsilon n)$. Similarly, $(b - a)^2 \tilde{S}^{2*}$ has location parameter $(b - a)^2 \tilde{S}^2 = S^2$ and scale parameter $(b - a)^2/(\varepsilon n)$. The result follows. $\square$

By this result, it is equivalent to release $\tilde{Y}^*$ and S^{2*} on the $[a, b]$ scale directly via the Laplace distribution or to re-scale the data to the $[0, 1]$ scale, release $\tilde{Y}^*$ and $\tilde{S}^{2*}$, and then use the relationships in Observation 3.1 to convert back to the $[a, b]$ scale.

APPENDIX B. THE TRUNCATED GAMMA MIXTURE DISTRIBUTION

This section provides additional details related to the truncated gamma mixture (abbreviated as TGM) distribution. The TGM distribution has four parameters: α is analogous to the gamma distribution's scale parameter, β is analogous to the gamma distribution's rate parameter, λ determines how far the TGM distribution is from a gamma distribution, and τ determines the point of truncation. When $\tau > 0$, the distribution is a mixture of $\mathbf{Gam}_{(0, \tau]}(\alpha, \beta - \lambda)$ and $\mathbf{Gam}_{(\tau, \infty)}(\alpha, \beta + \lambda)$ with the mixture weights in (3.3). When $\tau \leq 0$, the distribution is equivalent to $\mathbf{Gam}(\alpha, \beta + \lambda)$.

Figure 7 compares the probability density of a TGM distribution to the probability density of the gamma distribution with the same α and β . We see that the shapes are similar, but the rates are different on each side of τ . Notably, the distribution is continuous but is clearly not differentiable at τ .

Algorithm 1 provides a straightforward procedure for sampling from the TGM distribution. It assumes one is able to sample from a truncated gamma distribution, which can be done via standard tools such as the `rtrunc` function from the `truncdist` package in R (Novomestky and Nadarajah, 2016, License: GPL ≥ 2).

We now prove that Algorithm 1 correctly samples from the TGM distribution.

Theorem B.1. *X sampled via Algorithm 1 has distribution $X \sim \mathbf{TGM}(\alpha, \beta, \lambda, \tau)$.*

Proof. If $\tau \leq 0$, then $X \sim \mathbf{Gam}(\alpha, \beta + \lambda)$, as desired.

We now consider the case where $\tau > 0$. Let $F_1(x)$ and $f_1(x)$ be the CDF and PDF, respectively, of $\mathbf{Gam}(\alpha, \beta - \lambda)$ truncated to $(0, \tau]$. Let $F_2(x)$ and $f_2(x)$ be the CDF and

Algorithm 1: Sample $X \sim \text{TGM}(\alpha, \beta, \lambda, \tau)$

Input: $\alpha, \beta, \lambda, \tau$
if $\tau \leq 0$ **then**
 Sample $X \sim \text{Gam}(\alpha, \beta + \lambda)$
else
 Compute π_1 via (3.3)
 Sample $U \sim \text{Unif}(0, 1)$
 if $U \leq \pi_1$ **then**
 Sample $X \sim \text{Gam}(\alpha, \beta - \lambda)$ truncated to $(0, \tau]$
 else
 Sample $X \sim \text{Gam}(\alpha, \beta + \lambda)$ truncated to (τ, ∞) ;
Output: X

PDF, respectively, of $\text{Gam}(\alpha, \beta + \lambda)$ truncated to (τ, ∞) . By the law of total probability, the CDF of X is

$$P[X \leq x] = P[X \leq x \mid U \leq \pi_1] P[U \leq \pi_1] + P[X \leq x \mid U > \pi_1] P[U > \pi_1] \quad (\text{B.1})$$

$$= \pi_1 F_1(x) + \pi_2 F_2(x). \quad (\text{B.2})$$

This function is differentiable everywhere except $x = \tau$. When $x \neq \tau$, we may take the derivative with respect to x to obtain the PDF

$$p(x) = \frac{d}{dx} [\pi_1 F_1(x) + \pi_2 F_2(x)] \quad (\text{B.3})$$

$$= \pi_1 f_1(x) + \pi_2 f_2(x) \quad (\text{B.4})$$

$$= \pi_1 \frac{(\beta - \lambda)^\alpha}{\gamma(\alpha, (\beta - \lambda)\tau)} x^{\alpha-1} e^{-(\beta-\lambda)x} \mathbf{1}[x \leq \tau] + \pi_2 \frac{(\beta + \lambda)^\alpha}{\Gamma(\alpha, (\beta + \lambda)\tau)} x^{\alpha-1} e^{-(\beta+\lambda)x} \mathbf{1}[x > \tau]. \quad (\text{B.5})$$

Thus, we have shown that $X \sim \text{TGM}(\alpha, \beta, \lambda, \tau)$ for all $X \neq \tau$. Since the event $X = \tau$ occurs with probability zero, this completes the proof. $\square$

The TGM distribution primarily arises in the following setting.

Theorem B.2. *If $\tau \mid \mu \sim \text{Lap}(\mu, 1/\lambda)$ and $\mu \sim \text{Gam}(\alpha, \beta)$ for $\beta > \lambda$, then $\mu \mid \tau \sim \text{TGM}(\alpha, \beta, \lambda, \tau)$.*

Proof. By Bayes' Theorem, the desired distribution is

$$p(\mu \mid \tau) = \frac{p(\tau \mid \mu) p(\mu)}{\int_0^\infty p(\tau \mid \mu) p(\mu) d\mu}. \quad (\text{B.6})$$

We begin by investigating the numerator of (B.6), which has the following form.

$$p(\tau \mid \mu) p(\mu) = \frac{\lambda}{2} e^{-\lambda|\tau-\mu|} \frac{\beta^\alpha}{\Gamma(\alpha)} \mu^{\alpha-1} e^{-\beta\mu} \quad (\text{B.7})$$

$$= \frac{\lambda\beta^\alpha}{2\Gamma(\alpha)} \mu^{\alpha-1} \left[e^{-\lambda\tau} e^{-(\beta-\lambda)\mu} \mathbf{1}[\mu \leq \tau] + e^{\lambda\tau} e^{-(\beta+\lambda)\mu} \mathbf{1}[\mu > \tau] \right] \quad (\text{B.8})$$

$$= \frac{\lambda\beta^\alpha}{2\Gamma(\alpha)} \left[e^{-\lambda\tau} \mu^{\alpha-1} e^{-(\beta-\lambda)\mu} \mathbf{1}[\mu \leq \tau] + e^{\lambda\tau} \mu^{\alpha-1} e^{-(\beta+\lambda)\mu} \mathbf{1}[\mu > \tau] \right] \quad (\text{B.9})$$

$$= \frac{\lambda\beta^\alpha}{2\Gamma(\alpha)} \left[e^{-\lambda\tau} \frac{\gamma(\alpha, (\beta-\lambda)\tau)}{(\beta-\lambda)^\alpha} \frac{(\beta-\lambda)^\alpha}{\gamma(\alpha, (\beta-\lambda)\tau)} \mu^{\alpha-1} e^{-(\beta-\lambda)\mu} \mathbf{1}[\mu \leq \tau] \right. \\ \left. + e^{\lambda\tau} \frac{\Gamma(\alpha, (\beta+\lambda)\tau)}{(\beta+\lambda)^\alpha} \frac{(\beta+\lambda)^\alpha}{\Gamma(\alpha, (\beta+\lambda)\tau)} \mu^{\alpha-1} e^{-(\beta+\lambda)\mu} \mathbf{1}[\mu > \tau] \right]. \quad (\text{B.10})$$

The denominator of (B.6) is then as follows.

$$\int_0^\infty p(\tau \mid \mu) p(\mu) d\mu = \frac{\lambda\beta^\alpha}{2\Gamma(\alpha)} \left[e^{-\lambda\tau} \frac{\gamma(\alpha, (\beta-\lambda)\tau)}{(\beta-\lambda)^\alpha} \int_0^\tau \frac{(\beta-\lambda)^\alpha}{\gamma(\alpha, (\beta-\lambda)\tau)} \mu^{\alpha-1} e^{-(\beta-\lambda)\mu} d\mu \right. \\ \left. + e^{\lambda\tau} \frac{\Gamma(\alpha, (\beta+\lambda)\tau)}{(\beta+\lambda)^\alpha} \int_\tau^\infty \frac{(\beta+\lambda)^\alpha}{\Gamma(\alpha, (\beta+\lambda)\tau)} \mu^{\alpha-1} e^{-(\beta+\lambda)\mu} d\mu \right] \quad (\text{B.11})$$

$$= \frac{\lambda\beta^\alpha}{2\Gamma(\alpha)} \left[e^{-\lambda\tau} \frac{\gamma(\alpha, (\beta-\lambda)\tau)}{(\beta-\lambda)^\alpha} + e^{\lambda\tau} \frac{\Gamma(\alpha, (\beta+\lambda)\tau)}{(\beta+\lambda)^\alpha} \right]. \quad (\text{B.12})$$

The second equality follows by recognizing the integrands as the PDFs of truncated gamma distributions. Thus, the distribution of $\mu \mid \tau$ is

$$p(\mu \mid \tau) = \frac{e^{-\lambda\tau} \frac{\gamma(\alpha, (\beta-\lambda)\tau)}{(\beta-\lambda)^\alpha}}{e^{-\lambda\tau} \frac{\gamma(\alpha, (\beta-\lambda)\tau)}{(\beta-\lambda)^\alpha} + e^{\lambda\tau} \frac{\Gamma(\alpha, (\beta+\lambda)\tau)}{(\beta+\lambda)^\alpha}} \frac{(\beta-\lambda)^\alpha}{\gamma(\alpha, (\beta-\lambda)\tau)} \mu^{\alpha-1} e^{-(\beta-\lambda)\mu} \mathbf{1}[\mu \leq \tau] \\ + \frac{e^{\lambda\tau} \frac{\Gamma(\alpha, (\beta+\lambda)\tau)}{(\beta+\lambda)^\alpha}}{e^{-\lambda\tau} \frac{\gamma(\alpha, (\beta-\lambda)\tau)}{(\beta-\lambda)^\alpha} + e^{\lambda\tau} \frac{\Gamma(\alpha, (\beta+\lambda)\tau)}{(\beta+\lambda)^\alpha}} \frac{(\beta+\lambda)^\alpha}{\Gamma(\alpha, (\beta+\lambda)\tau)} \mu^{\alpha-1} e^{-(\beta+\lambda)\mu} \mathbf{1}[\mu > \tau], \quad (\text{B.13})$$

which we recognize as the PDF of a **TGM**($\alpha, \beta, \lambda, \tau$). $\square$

APPENDIX C. ALGORITHMS AND FULL CONDITIONALS

In this section, we provide explicit algorithms for implementing the Gibbs samplers described in the main body. In the algorithms, we denote the t^{th} draw from the posterior for a variable x as $x_{(t)}$.

Algorithm 2: Run our proposed Gibbs sampler.

Input: Data $(\bar{Y}^*, S^{2*})$, privacy budget $(\varepsilon_1, \varepsilon_2)$, sample size n , iterations T , hyperparameters $(\mu_0, \sigma_0^2, \kappa_0, \nu_0)$
Initialize $\sigma_{(0)}^2$, $\bar{Y}_{(0)}$, $S_{(0)}^2$, and $\omega_{(0)}^2$;
for t **in** $1 : T$ **do**
 Sample $\mu_{(t)}$ from $(\mu \mid \sigma^2 = \sigma_{(t-1)}^2, \omega^2 = \omega_{(t-1)}^2, \bar{Y} = \bar{Y}_{(t-1)}, S^2 = S_{(t-1)}^2, \bar{Y}^*, S^{2*})$
 Sample $\sigma_{(t)}^2$ from $(\sigma^2 \mid \mu = \mu_{(t)}, \omega^2 = \omega_{(t-1)}^2, \bar{Y} = \bar{Y}_{(t-1)}, S^2 = S_{(t-1)}^2, \bar{Y}^*, S^{2*})$
 Sample $\bar{Y}_{(t)}$ from $(\bar{Y} \mid \mu = \mu_{(t)}, \sigma^2 = \sigma_{(t)}^2, \omega^2 = \omega_{(t-1)}^2, S^2 = S_{(t-1)}^2, \bar{Y}^*, S^{2*})$
 Sample $S_{(t)}^2$ from $(S^2 \mid \mu = \mu_{(t)}, \sigma^2 = \sigma_{(t)}^2, \omega^2 = \omega_{(t-1)}^2, \bar{Y} = \bar{Y}_{(t)}, \bar{Y}^*, S^{2*})$
 Sample $\omega_{(t)}^2$ from $(\omega^2 \mid \mu = \mu_{(t)}, \sigma^2 = \sigma_{(t)}^2, \bar{Y} = \bar{Y}_{(t)}, S^2 = S_{(t)}^2, \bar{Y}^*, S^{2*})$
Output: $(\mu_{(1)}, \dots, \mu_{(T)}), (\sigma_{(1)}^2, \dots, \sigma_{(T)}^2)$

C.1. Algorithm for Sampling Gaussian Model Parameters Via a Measurement Error Model. In Algorithm 2, the first step is to initialize values for each of the variables, which we describe further in the following paragraph. The analyst then repeatedly samples from the full conditionals for each of the variables. The forms of the full conditionals under different assumptions are described in Appendices C.1.1—C.1.4. The final step is to output the posterior draws for the two parameters.

The posterior distribution is not very sensitive to the initial values, but well-thought-out values may reduce the runtime. In our examples, we set $\sigma_{(0)}^2 = S_{(0)}^2 = S^{2*}$ if $S^{2*} \in (0, 1/4)$ and $\bar{Y}_{(0)} = \bar{Y}^*$ if $\bar{Y}^* \in [0, 1]$. If $\bar{Y}^*$ or S^{2*} are not in these regions, they are set near the closest boundary (although note that $\sigma_{(0)}^2$ and $S_{(0)}^2$ should not be set exactly equal to zero). We set $\omega_{(0)}^2 = 2/(\varepsilon_1^2 n^2)$.

We now summarize the full conditional under different prior distributions for the DP univariate Gaussian setting with $\bar{Y}^* \sim \text{Lap}(\bar{Y}, 1/(\varepsilon_1 n))$ and $S^{2*} \sim \text{Lap}(S^2, 1/(\varepsilon_2 n))$. Subscripts on a distribution denote truncation of the distribution to an interval. The set A is as defined in (3.4).

C.1.1. Conjugate Prior Without Constraints. If the prior is $\sigma^2 \sim \text{IG}(\nu_0/2, \nu_0 \sigma_0^2/2)$ and $\mu \mid \sigma^2 \sim \mathcal{N}(\mu_0, \sigma^2/\kappa_0)$ with no constraints enforced, then

$$\mu \mid \sigma^2, \omega^2, \bar{Y}, S^2, \bar{Y}^*, S^{2*} \sim \mathcal{N}\left(\frac{n\bar{Y} + \kappa_0 \mu_0}{n + \kappa_0}, \frac{\sigma^2}{n + \kappa_0}\right) \quad (\text{C.1})$$

$$\sigma^2 \mid \mu, \omega^2, \bar{Y}, S^2, \bar{Y}^*, S^{2*} \sim \text{IG}_{(0, \frac{n-1}{2n\varepsilon_2})}\left(\frac{n + \nu_0}{2}, \frac{\nu_0 \sigma_0^2 + (n-1)S^2 + n(\bar{Y} - \mu)^2}{2}\right) \quad (\text{C.2})$$

$$\bar{Y} \mid \mu, \sigma^2, \omega^2, S^2, \bar{Y}^*, S^{2*} \sim \mathcal{N}\left(\frac{\frac{\bar{Y}^*}{\omega^2} + \frac{n\mu}{\sigma^2}}{\frac{1}{\omega^2} + \frac{n}{\sigma^2}}, \frac{1}{\frac{1}{\omega^2} + \frac{n}{\sigma^2}}\right) \quad (\text{C.3})$$

$$1/\omega^2 \mid \mu, \sigma^2, \bar{Y}, S^2, \bar{Y}^*, S^{2*} \sim \text{InvGaus}\left(\frac{\varepsilon_1 n}{|\bar{Y}^* - \bar{Y}|}, \varepsilon_1^2 n^2\right) \quad (\text{C.4})$$

$$S^2 \mid \mu, \sigma^2, \omega^2, \bar{Y}, \bar{Y}^*, S^{2*} \sim \text{TGM}\left(\frac{n-1}{2}, \frac{n-1}{2\sigma^2}, n\varepsilon_2, S^{2*}\right). \quad (\text{C.5})$$

C.1.2. Truncated Conjugate Prior To Enforce Constraints. If the prior is $\sigma^2 \sim \text{IG}(\nu_0/2, \nu_0\sigma_0^2/2)$ and $\mu \mid \sigma^2 \sim \mathcal{N}(\mu_0, \sigma^2/\kappa_0)$ constrained such that $(\mu, \sigma^2) \in A$, then

$$\mu \mid \sigma^2, \omega^2, \bar{Y}, S^2, \bar{Y}^*, S^{2*} \sim \mathcal{N}_{\left[1/2 - \sqrt{1/4 - \sigma^2}, 1/2 + \sqrt{1/4 - \sigma^2}\right]} \left(\frac{n\bar{Y} + \kappa_0\mu_0}{n + \kappa_0}, \frac{\sigma^2}{n + \kappa_0} \right) \quad (\text{C.6})$$

$$\sigma^2 \mid \mu, \omega^2, \bar{Y}, S^2, \bar{Y}^*, S^{2*} \sim \text{IG}_{\left[0, \min\{\mu(1-\mu), \frac{n-1}{2n\varepsilon_2}\}\right]} \left(\frac{n + \nu_0}{2}, \frac{\nu_0\sigma_0^2 + (n-1)S^2 + n(\bar{Y} - \mu)^2}{2} \right) \quad (\text{C.7})$$

$$\bar{Y} \mid \mu, \sigma^2, \omega^2, S^2, \bar{Y}^*, S^{2*} \sim \mathcal{N}_{\left[1/2 - \sqrt{1/4 - \frac{n-1}{n}S^2}, 1/2 + \sqrt{1/4 - \frac{n-1}{n}S^2}\right]} \left(\frac{\frac{\bar{Y}^*}{\omega^2} + \frac{n\mu}{\sigma^2}}{\frac{1}{\omega^2} + \frac{n}{\sigma^2}}, \frac{1}{\frac{1}{\omega^2} + \frac{n}{\sigma^2}} \right) \quad (\text{C.8})$$

$$1/\omega^2 \mid \mu, \sigma^2, \bar{Y}, S^2, \bar{Y}^*, S^{2*} \sim \text{InvGaus} \left(\frac{\varepsilon_1 n}{|\bar{Y}^* - \bar{Y}|}, \varepsilon_1^2 n^2 \right) \quad (\text{C.9})$$

$$S^2 \mid \mu, \sigma^2, \omega^2, \bar{Y}, \bar{Y}^*, S^{2*} \sim \text{TGM}_{\left(0, \frac{n-1}{n-1}\bar{Y}(1-\bar{Y})\right]} \left(\frac{n-1}{2}, \frac{n-1}{2\sigma^2}, n\varepsilon_2, S^{2*} \right). \quad (\text{C.10})$$

C.1.3. Improper Prior Without Constraints. : If the prior is $p(\mu, \sigma^2) \propto 1$ with no constraints enforced, then

$$\mu \mid \sigma^2, \omega^2, \bar{Y}, S^2, \bar{Y}^*, S^{2*} \sim \mathcal{N} \left(\bar{Y}, \frac{\sigma^2}{n} \right) \quad (\text{C.11})$$

$$\sigma^2 \mid \mu, \omega^2, \bar{Y}, S^2, \bar{Y}^*, S^{2*} \sim \text{IG}_{\left(0, \frac{n-1}{2n\varepsilon_2}\right)} \left(\frac{n-2}{2}, \frac{(n-1)S^2 + n(\bar{Y} - \mu)^2}{2} \right) \quad (\text{C.12})$$

$$\bar{Y} \mid \mu, \sigma^2, \omega^2, S^2, \bar{Y}^*, S^{2*} \sim \mathcal{N} \left(\frac{\frac{\bar{Y}^*}{\omega^2} + \frac{n\mu}{\sigma^2}}{\frac{1}{\omega^2} + \frac{n}{\sigma^2}}, \frac{1}{\frac{1}{\omega^2} + \frac{n}{\sigma^2}} \right) \quad (\text{C.13})$$

$$1/\omega^2 \mid \mu, \sigma^2, \bar{Y}, S^2, \bar{Y}^*, S^{2*} \sim \text{InvGaus} \left(\frac{\varepsilon_1 n}{|\bar{Y}^* - \bar{Y}|}, \varepsilon_1^2 n^2 \right) \quad (\text{C.14})$$

$$S^2 \mid \mu, \sigma^2, \omega^2, \bar{Y}, \bar{Y}^*, S^{2*} \sim \text{TGM} \left(\frac{n-1}{2}, \frac{n-1}{2\sigma^2}, n\varepsilon_2, S^{2*} \right). \quad (\text{C.15})$$

C.1.4. Uniform Prior Over Feasible Region. If the prior is uniform over the set A defined in (3.4), then

$$\mu \mid \sigma^2, \omega^2, \bar{Y}, S^2, \bar{Y}^*, S^{2*} \sim \mathcal{N}_{\left[1/2 - \sqrt{1/4 - \sigma^2}, 1/2 + \sqrt{1/4 - \sigma^2}\right]} \left(\bar{Y}, \frac{\sigma^2}{n} \right) \quad (\text{C.16})$$

$$\sigma^2 \mid \mu, \omega^2, \bar{Y}, S^2, \bar{Y}^*, S^{2*} \sim \text{IG}_{\left[0, \min\{\mu(1-\mu), \frac{n-1}{2n\varepsilon_2}\right]} \left(\frac{n-2}{2}, \frac{(n-1)S^2 + n(\bar{Y} - \mu)^2}{2} \right) \quad (\text{C.17})$$

$$\bar{Y} \mid \mu, \sigma^2, \omega^2, S^2, \bar{Y}^*, S^{2*} \sim \mathcal{N}_{\left[1/2 - \sqrt{1/4 - \frac{n-1}{n}S^2}, 1/2 + \sqrt{1/4 - \frac{n-1}{n}S^2}\right]} \left(\frac{\frac{\bar{Y}^*}{\omega^2} + \frac{n\mu}{\sigma^2}}{\frac{1}{\omega^2} + \frac{n}{\sigma^2}}, \frac{1}{\frac{1}{\omega^2} + \frac{n}{\sigma^2}} \right) \quad (\text{C.18})$$

$$1/\omega^2 \mid \mu, \sigma^2, \bar{Y}, S^2, \bar{Y}^*, S^{2*} \sim \text{InvGaus} \left(\frac{\varepsilon_1 n}{|\bar{Y}^* - \bar{Y}|}, \varepsilon_1^2 n^2 \right) \quad (\text{C.19})$$

$$S^2 \mid \mu, \sigma^2, \omega^2, \bar{Y}, \bar{Y}^*, S^{2*} \sim \text{TGM}_{(0, \frac{n}{n-1}\bar{Y}(1-\bar{Y}))} \left(\frac{n-1}{2}, \frac{n-1}{2\sigma^2}, n\varepsilon_2, S^{2*} \right). \quad (\text{C.20})$$

C.2. Derivation of Full Conditionals for the Measurement Error Model. We now derive the full conditionals for the unconstrained measurement error model with the informative prior. This yields the full conditionals in equations (C.1)-(C.5). The derivations for the full conditionals in the other settings are similar and described in Appendices C.1.1—C.1.4. See Appendix C.3 for discussion of how the full conditionals are modified in the presence of constraints.

Recall that if $Y_1, \dots, Y_n \stackrel{iid}{\sim} \mathcal{N}(\mu, \sigma^2)$, then the distributions of the sufficient statistics are

$$\bar{Y} \mid \mu, \sigma^2 \sim \mathcal{N} \left(\mu, \frac{\sigma^2}{n} \right), \quad S^2 \mid \sigma^2 \sim \text{Gam} \left(\frac{n-1}{2}, \frac{n-1}{2\sigma^2} \right), \quad (\text{C.21})$$

where $(\bar{Y} \perp\!\!\!\perp S^2) \mid \mu, \sigma^2$. Assuming that the $Y_i \in [0, 1]$, then the statistics released via the Laplace mechanism have distributions

$$\bar{Y}^* \mid \bar{Y} \sim \text{Lap} \left(\bar{Y}, \frac{1}{\varepsilon_1 n} \right), \quad S^{2*} \mid S^2 \sim \text{Lap} \left(S^2, \frac{1}{\varepsilon_2 n} \right). \quad (\text{C.22})$$

The release mechanism is such that $(\bar{Y}^* \perp\!\!\!\perp S^{2*}) \mid \bar{Y}, S^2$. It is convenient to replace the above distribution of $\bar{Y}^*$ with the following equivalent formulation, as proposed by Park and Casella (2008) and used in Bernstein and Sheldon (2018, 2019).

$$\bar{Y}^* \mid \bar{Y}, \omega^2 \sim \mathcal{N}(\bar{Y}, \omega^2), \quad \omega^2 \sim \text{Exp} \left(\frac{\varepsilon_1^2 n^2}{2} \right). \quad (\text{C.23})$$

Finally, we take the conjugate prior from the public setting

$$\sigma^2 \sim \text{IG} \left(\frac{\nu_0}{2}, \frac{\nu_0 \sigma_0^2}{2} \right), \quad \mu \mid \sigma^2 \sim \mathcal{N} \left(\mu_0, \frac{\sigma^2}{\kappa_0} \right). \quad (\text{C.24})$$

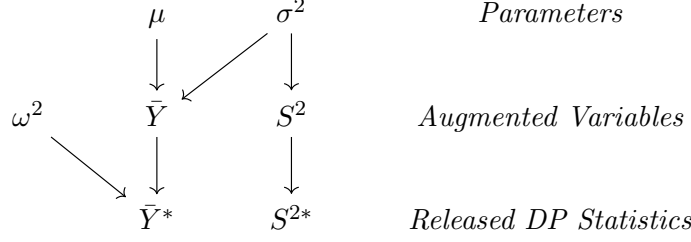


Figure 8: A graphical model representing the structure of the univariate Gaussian setting.

Figure 8 presents the graphical model corresponding to this likelihood, which we can factor as follows.

$$p(\bar{Y}^*, S^{2*}, \bar{Y}, S^2, \omega^2 \mid \mu, \sigma^2) = p(\bar{Y}^* \mid \bar{Y}, \omega^2) p(S^{2*} \mid S^2) p(\bar{Y} \mid \mu, \sigma^2) p(S^2 \mid \sigma^2) p(\omega^2) \quad (\text{C.25})$$

$$= \mathcal{N}(\bar{Y}^*; \bar{Y}, \omega^2) \text{Lap}\left(S^{2*}; S^2, \frac{1}{\varepsilon_2 n}\right) \mathcal{N}\left(\bar{Y}; \mu, \frac{\sigma^2}{n}\right) \text{Gam}\left(S^2; \frac{n-1}{2}, \frac{n-1}{2\sigma^2}\right) \text{Exp}\left(\omega^2; \frac{\varepsilon_1^2 n^2}{2}\right). \quad (\text{C.26})$$

We now examine each full conditional. The full conditional for μ is

$$p(\mu \mid \bar{Y}^*, S^{2*}, \bar{Y}, S^2, \omega^2, \sigma^2) \propto p(\bar{Y} \mid \mu, \sigma^2) p(\mu \mid \sigma^2) = \mathcal{N}\left(\bar{Y}; \mu, \frac{\sigma^2}{n}\right) \mathcal{N}\left(\mu; \mu_0, \frac{\sigma^2}{\kappa_0}\right), \quad (\text{C.27})$$

which we recognize as the full conditional from the public setting, as derived in Chapter 5 of Hoff (2009). Similarly, the full conditional for σ^2 is

$$p(\sigma^2 \mid \bar{Y}^*, S^{2*}, \bar{Y}, S^2, \omega^2, \mu) \propto p(\bar{Y} \mid \mu, \sigma^2) p(S^2 \mid \sigma^2) p(\sigma^2) \quad (\text{C.28})$$

$$= \mathcal{N}\left(\bar{Y}; \mu, \frac{\sigma^2}{n}\right) \text{Gam}\left(S^2; \frac{n-1}{2}, \frac{n-1}{2\sigma^2}\right) \text{IG}\left(\frac{\nu_0}{2}, \frac{\nu_0 \sigma_0^2}{2}\right). \quad (\text{C.29})$$

This is the same full conditional as in the public setting, as derived in Chapter 6 of Hoff (2009). Thus, the full conditionals have the form

$$\mu \mid \sigma^2, \omega^2, \bar{Y}, S^2, \bar{Y}^*, S^{2*} \sim \mathcal{N}\left(\frac{n\bar{Y} + \kappa_0 \mu_0}{n + \kappa_0}, \frac{\sigma^2}{n + \kappa_0}\right) \quad (\text{C.30})$$

$$\sigma^2 \mid \mu, \omega^2, \bar{Y}, S^2, \bar{Y}^*, S^{2*} \sim \text{IG}\left(\frac{n + \nu_0}{2}, \frac{\nu_0 \sigma_0^2 + (n-1)S^2 + n(\bar{Y} - \mu)^2}{2}\right). \quad (\text{C.31})$$

We next examine the full conditional for $\bar{Y}$, which is

$$p(\bar{Y} \mid \bar{Y}^*, S^{2*}, S^2, \omega^2, \mu, \sigma^2) \propto p(\bar{Y}^* \mid \bar{Y}, \omega^2) p(\bar{Y} \mid \mu, \sigma^2) = \mathcal{N}(\bar{Y}^*; \bar{Y}, \omega^2) \mathcal{N}\left(\bar{Y}; \mu, \frac{\sigma^2}{n}\right). \quad (\text{C.32})$$

We recognize this is equivalent to a Gaussian model with known variance and a Gaussian prior on the mean. By Chapter 5 of [Hoff \(2009\)](#), the full conditional for $\bar{Y}$ is

$$\bar{Y} \mid \mu, \sigma^2, \omega^2, S^2, \bar{Y}^*, S^{2*} \sim \mathcal{N} \left(\frac{\bar{Y}^* + \frac{n\mu}{\omega^2}}{\frac{1}{\omega^2} + \frac{n}{\sigma^2}}, \frac{1}{\frac{1}{\omega^2} + \frac{n}{\sigma^2}} \right). \quad (\text{C.33})$$

The full conditional for ω^2 is as follows, expressed in terms of its inverse, $1/\omega^2$.

$$p(1/\omega^2 \mid \bar{Y}^*, S^{2*}, \bar{Y}, S^2, \mu, \sigma^2) \propto p(\bar{Y}^* \mid \bar{Y}, 1/\omega^2) p(1/\omega^2) \quad (\text{C.34})$$

$$= \mathcal{N}(\bar{Y}^*; \bar{Y}, (1/\omega^2)^{-1}) \text{IG} \left(1/\omega^2; 1, \frac{\varepsilon_1^2 n^2}{2} \right). \quad (\text{C.35})$$

We recognize this form from the Bayesian LASSO ([Park and Casella, 2008](#)), which yields the distribution

$$1/\omega^2 \mid \mu, \sigma^2, \bar{Y}, S^2, \bar{Y}^*, S^{2*} \sim \text{InvGaus} \left(\frac{\varepsilon_1 n}{|\bar{Y}^* - \bar{Y}|}, \varepsilon_1^2 n^2 \right). \quad (\text{C.36})$$

Finally, the full conditional for S^2 is

$$p(S^2 \mid \bar{Y}^*, S^{2*}, \bar{Y}, \omega^2, \mu, \sigma^2) \propto p(S^{2*} \mid S^2) p(S^2 \mid \sigma^2) \quad (\text{C.37})$$

$$= \text{Lap} \left(S^{2*}; S^2, \frac{1}{\varepsilon_2 n} \right) \text{Gam} \left(S^2; \frac{n-1}{2}, \frac{n-1}{2\sigma^2} \right). \quad (\text{C.38})$$

By Theorem [B.2](#) when $(n-1)/(2\sigma^2) > \varepsilon_2 n$, the distribution is

$$S^2 \mid \mu, \sigma^2, \omega^2, \bar{Y}, \bar{Y}^*, S^{2*} \sim \text{TGM} \left(\frac{n-1}{2}, \frac{n-1}{2\sigma^2}, n\varepsilon_2, S^{2*} \right). \quad (\text{C.39})$$

Note that by Observation [3.4](#), which is repeated in Appendix [D](#) below, since $Y_i \in [0, 1]$ it follows that $\sigma^2 \leq 1/4$. Thus, we may use this full conditional when

$$\varepsilon_2 < \frac{n-1}{n} \cdot \frac{1}{2\sigma^2} \leq 2 \frac{n-1}{n}. \quad (\text{C.40})$$

If $\varepsilon_2 > 2(n-1)/n$, then the full conditional for S^2 is not necessarily a TGM without additional assumptions. If reasonable for the application, the analyst could impose the additional constraint that $\sigma^2 < (n-1)/(n\varepsilon_2)$, which is sufficient to ensure the full conditional for S^2 is of the form in [\(C.39\)](#). Alternatively, if n is large, an analyst could use an approximation of the full conditional for S^2 , such as the proposals of [Bernstein and Sheldon \(2018, 2019\)](#).

C.3. Derivation of Full Conditionals with Constraints. We now demonstrate how the full conditional distributions in Appendix [C.2](#) change when constraints are enforced on the prior distribution and sufficient statistics. We present two example derivations for the measurement error model to demonstrate the ideas. The other full conditionals with constraints that are enumerated in Appendix [C.1](#) can be derived analogously.

As in Appendix [C.2](#), let $p(\mu, \sigma^2) = p(\mu \mid \sigma^2) p(\sigma^2)$ be the joint prior distribution without constraints. Define $\tilde{p}(\mu, \sigma^2)$ to be the joint prior distribution with the constraint that $(\mu, \sigma^2) \in A$. That is

$$\tilde{p}(\mu, \sigma^2) = c_1 p(\mu, \sigma^2) \mathbf{1}[(\mu, \sigma^2) \in A] \quad (\text{C.41})$$

$$= c_1 p(\mu \mid \sigma^2) p(\sigma^2) \mathbf{1}[(\mu, \sigma^2) \in A], \quad (\text{C.42})$$

where c_1 is a normalizing constant that is not a function of μ or σ^2 . Similarly, let $p(\bar{Y}, S^2 \mid \mu, \sigma^2) = p(\bar{Y} \mid \mu, \sigma^2)p(S^2 \mid \sigma^2)$ be the likelihood for the sufficient statistics without constraints. Let B be the constrained region for the likelihood implied by Observation 3.5, i.e.,

$$B = \{(\bar{Y}, S^2) : \bar{Y} \in [0, 1] \text{ and } S^2 \in [0, n/(n-1)\bar{Y}(1-\bar{Y})]\} \quad (\text{C.43})$$

Define $\tilde{p}(\bar{Y}, S^2 \mid \mu, \sigma^2)$ to be the likelihood with the constraint that $(\bar{Y}, S^2) \in B$. That is,

$$\tilde{p}(\bar{Y}, S^2 \mid \mu, \sigma^2) = c_2 p(\bar{Y}, S^2 \mid \mu, \sigma^2) \mathbf{1}[(\bar{Y}, S^2) \in B] \quad (\text{C.44})$$

$$= c_2 p(\bar{Y} \mid \mu, \sigma^2) p(S^2 \mid \sigma^2) \mathbf{1}[(\bar{Y}, S^2) \in B], \quad (\text{C.45})$$

where c_2 is a normalizing constant that is not a function of $\bar{Y}$ or S^2 .

Applying (C.42) and (C.45) along with the unconstrained posterior from (C.30), yields the following full conditional for μ .

$$\tilde{p}(\mu \mid \bar{Y}^*, S^{2*}, \bar{Y}, S^2, \omega^2, \sigma^2) \propto \tilde{p}(\bar{Y}, S^2 \mid \mu, \sigma^2) \tilde{p}(\mu, \sigma^2) \quad (\text{C.46})$$

$$\propto p(\bar{Y} \mid \mu, \sigma^2) p(\mu \mid \sigma^2) \mathbf{1}[(\mu, \sigma^2) \in A] \quad (\text{C.47})$$

$$\propto \mathcal{N}\left(\bar{Y}; \mu, \frac{\sigma^2}{n}\right) \mathcal{N}\left(\mu; \mu_0, \frac{\sigma^2}{\kappa_0}\right) \mathbf{1}[(\mu, \sigma^2) \in A] \quad (\text{C.48})$$

$$\propto \mathcal{N}\left(\mu; \frac{n\bar{Y} + \kappa_0\mu_0}{n + \kappa_0}, \frac{\sigma^2}{n + \kappa_0}\right) \mathbf{1}[(\mu, \sigma^2) \in A]. \quad (\text{C.49})$$

Combining this result with Observation 3.3 yields the form of the full conditional for μ in (C.11).

Similar derivations can be performed for the statistics' full conditionals with constraints. To demonstrate, consider the full conditional for $\bar{Y}$. Applying (C.45) and the unconstrained posterior from (C.33) yields

$$\tilde{p}(\bar{Y} \mid \bar{Y}^*, S^{2*}, S^2, \omega^2, \mu, \sigma^2) \propto p(\bar{Y}^* \mid \bar{Y}, \omega^2) \tilde{p}(\bar{Y}, S^2 \mid \mu, \sigma^2) \quad (\text{C.50})$$

$$\propto p(\bar{Y}^* \mid \bar{Y}, \omega^2) p(\bar{Y} \mid \mu, \sigma^2) \mathbf{1}[(\bar{Y}, S^2) \in B] \quad (\text{C.51})$$

$$\propto \mathcal{N}(\bar{Y}^*; \bar{Y}, \omega^2) \mathcal{N}\left(\bar{Y}; \mu, \frac{\sigma^2}{n}\right) \mathbf{1}[(\bar{Y}, S^2) \in B] \quad (\text{C.52})$$

$$\propto \mathcal{N}\left(\bar{Y}; \frac{\bar{Y}^* + \frac{n\mu}{\omega^2}}{\frac{1}{\omega^2} + \frac{n}{\sigma^2}}, \frac{1}{\frac{1}{\omega^2} + \frac{n}{\sigma^2}}\right) \mathbf{1}[(\bar{Y}, S^2) \in B]. \quad (\text{C.53})$$

Combining this result with Observation 3.5 yields the form of the full conditional for $\bar{Y}$ in (C.13).

APPENDIX D. CONSTRAINTS FOR THE UNIVARIATE GAUSSIAN SETTING

In this section, we provide proofs for results in Section 3.2 of the main body regarding constraints on parameters and statistics in the univariate Gaussian setting. The result regarding constraints on the parameters μ and σ^2 is restated below.

Observation 3.3. *Let $Y_i \in [0, 1]$ have moments $E[Y_i] = \mu$ and $V[Y_i] = \sigma^2$. It follows that $\sigma^2 \in [0, \mu(1-\mu)]$ and $\mu \in [1/2 - \sqrt{1/4 - \sigma^2}, 1/2 + \sqrt{1/4 - \sigma^2}]$.*

Proof. Since $Y_i \in [0, 1]$, it follows that $Y_i^2 \leq Y_i$. and so by monotonicity of expectation,

$$\sigma^2 = \text{Var}[Y_i] = E[Y_i^2] - E[Y_i]^2 \leq E[Y_i] - E[Y_i]^2 = \mu - \mu^2. \quad (\text{D.1})$$

Thus, since Y_i must have non-negative variance, if μ is known then $\sigma^2 \in [0, \mu(1 - \mu)]$. On the other hand, if σ^2 is known then by (D.1), μ must be such that $\mu^2 - \mu + \sigma^2 \leq 0$. Applying the quadratic formula, we find that $\mu^2 - \mu + \sigma^2$ has roots $\mu = 1/2 \pm \sqrt{1/4 - \sigma^2}$. Recognizing that $\mu^2 - \mu + \sigma^2 \leq 0$ when μ is between these roots yields the desired result. $\square$

Observation 3.4 is restated below.

Observation 3.4. *If $Y_i \in [0, 1]$, then $V[Y_i] = \sigma^2 \leq 1/4$.*

Proof. By Observation 3.3 $\sigma^2 \leq \mu(1 - \mu)$ and since $f(\mu) = \mu(1 - \mu)$ attains its maximum value at $\mu = 1/2$, it follows that $f(\mu) \leq f(1/2) = 1/4$ for all μ . The result follows. $\square$

The result regarding constraints on the sufficient statistics $\bar{Y}$ and S^2 is restated below.

Observation 3.5. *Let $Y_1, \dots, Y_n$ be such that each $Y_i \in [0, 1]$. Then $S^2 \in [0, n/(n-1)\bar{Y}(1 - \bar{Y})]$ and $\bar{Y} \in [1/2 - \sqrt{1/4 - (n-1)/n \cdot S^2}, 1/2 + \sqrt{1/4 - (n-1)/n \cdot S^2}]$.*

Proof. We begin by expressing S^2 in a more convenient form, as follows.

$$(n-1)S^2 = \sum_{i=1}^n (Y_i - \bar{Y})^2 = \sum_{i=1}^n Y_i^2 + n\bar{Y}^2 - 2\bar{Y} \sum_{i=1}^n Y_i = \sum_{i=1}^n Y_i^2 - n\bar{Y}^2. \quad (\text{D.2})$$

Since $Y_i \in [0, 1]$, we have that $Y_i \leq Y_i^2$ and so

$$(n-1)S^2 = \sum_{i=1}^n Y_i^2 - n\bar{Y}^2 \leq \sum_{i=1}^n Y_i - n\bar{Y}^2 = n\bar{Y}(1 - \bar{Y}). \quad (\text{D.3})$$

Since $(n-1)S^2$ is a sum of squares, it is non-negative and so if $\bar{Y}$ is known, $S^2 \in [0, n/(n-1)\bar{Y}(1 - \bar{Y})]$. On the other hand, if S^2 is known then $\bar{Y}$ must be such that $\bar{Y}^2 - \bar{Y} + (n-1)/n \cdot S^2 \leq 0$. By an analogous argument to the proof of Observation 3.3, it follows that $\bar{Y} \in [1/2 - \sqrt{1/4 - (n-1)/n \cdot S^2}, 1/2 + \sqrt{1/4 - (n-1)/n \cdot S^2}]$. $\square$

APPENDIX E. TRUNCATED VS CONSTRAINED POSTERIOR

In this section, we empirically demonstrate that the constrained posterior proposed in Section 3 is not merely a truncated version of the unconstrained posterior. We utilize the following example for the demonstration.

Example E.1. Let $Y_1, \dots, Y_n \sim \mathcal{N}(\mu, \sigma^2)$ where each $Y_i \in [0, 1]$ and $n = 25$. Suppose differentially private statistics $\bar{Y}^* = 1$ and $\bar{S}^{2*} = 0.05$ result from Laplace mechanisms using $\varepsilon_1 = \varepsilon_2 = 0.5$.

Figure 9 presents two versions of the analysis with the default prior distribution proposed in Section 4. The first version ignores the constraints (resulting in the draws displayed in the left panel of Figure 9) and then throws out the draws outside the feasible region (resulting in the draws in the middle panel of Figure 9). The second version incorporates the constraints, producing the posterior distribution in the right panel of Figure 9. To ensure differences in the distributions are not Monte Carlo noise, we examine the samplers after 1 million Gibbs

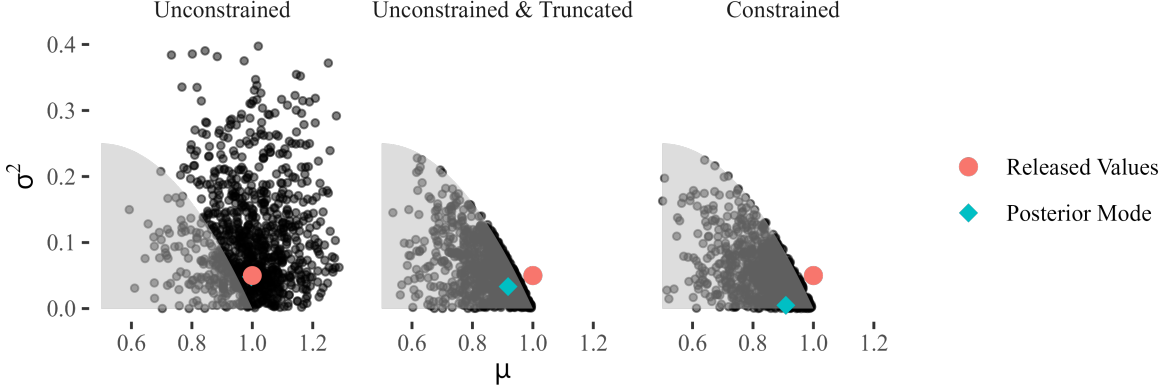


Figure 9: Posterior distribution for (μ, σ^2) in Example E.1 under default priors. The point $(\bar{Y}^* = 1, S^2 = 0.05)$ is represented by the red circle, and the analyst's posterior mode is represented by the blue diamond. The left panel displays draws from the sampler when constraints are not accounted for; the middle panel truncates these draws to the feasible region; and, the right panel displays draws with the constraints accounted for. The shaded area represents A , the feasible region for (μ, σ^2) . This plot is based on a subsample of 1,000 Gibbs iterations.

iterations. The truncation removes 70.0% of samples, resulting in approximately 300,000 remaining to produce point and interval estimates.

As evident in Figure 9, the distributions of σ^2 are different. In the constrained analysis, the posterior mode is $\hat{\sigma}^2 = 0.070^2$ with a 95% HPD interval $[0.001^2, 0.358^2]$. In the unconstrained-with-truncation analysis, the posterior mode is $\hat{\sigma}^2 = 0.182^2$ with a 95% HPD interval $[0.004^2, 0.368^2]$. The distribution of μ also differs although relatively less so. In the constrained analysis, the posterior mode is $\hat{\mu} = 0.909$ with 95% HPD credible interval $[0.653, 0.995]$. In the unconstrained-with-truncation analysis, the posterior mode is $\hat{\mu} = 0.918$ with 95% HPD interval $[0.670, 0.996]$.

We note that the differences in the two distributions can be harder to distinguish when using an informative prior distribution. For example, Figure 10 presents the corresponding analysis for the prior distribution in (3.1) with hyperparameters $\mu_0 = 0.8$, $\sigma_0^2 = 0.25^2$, $\kappa_0 = 0.5$, and $\nu_0 = 0.5$. We again utilize 1 million Gibbs iterations to ensure any differences are unlikely to be a result of Monte Carlo error. In the constrained analysis, the posterior modes are $\hat{\mu} = 0.914$ and $\hat{\sigma}^2 = 0.114^2$ with HPD 95% credible intervals $[0.695, 0.987]$ and $[0.054^2, 0.305^2]$, respectively. In the unconstrained analysis, the posterior modes are $\hat{\mu} = 0.930$ and $\hat{\sigma}^2 = 0.113^2$ with HPD 95% credible intervals $[0.717, 0.991]$ and $[0.053^2, 0.309^2]$, respectively.

APPENDIX F. DETERMINING WHETHER THE POSTERIOR DISTRIBUTION IS PROPER

In this section, we prove that the posterior distribution in the differentially private univariate Gaussian setting is not a proper probability distribution when $p(\mu, \sigma^2) \propto (\sigma^2)^{-1}$, but is a proper probability distribution when $p(\mu, \sigma^2) \propto 1$. The results are reproduced below.

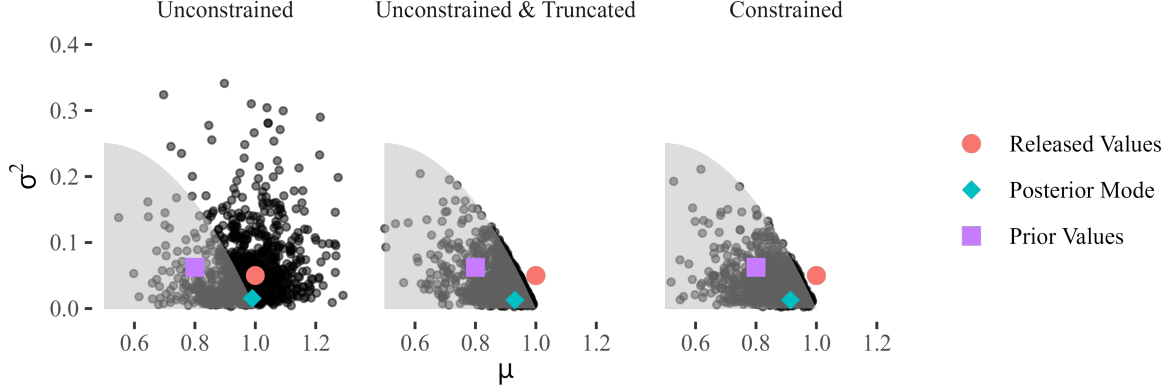


Figure 10: Posterior distribution for (μ, σ^2) in Example E.1 under an informative prior. The point $(\bar{Y}^* = 1, S^2 = 0.05)$ is represented by the red circle, the analyst's posterior mode is represented by the blue diamond, and the analyst's prior estimates $(\mu_0 = 0.8, \sigma^2 = 0.25^2)$ are represented by the purple square. The left panel provides the posterior when constraints are not accounted for; the middle panel provides the corresponding analysis truncated to the feasible region; and, the right panel presents the analysis with the constraints accounted for. The shaded area represents A , the feasible region for (μ, σ^2) . This plot is based on a subsample of 1,000 Gibbs iterations.

Theorem 4.1. For confidential data $Y_1, \dots, Y_n \stackrel{iid}{\sim} \mathcal{N}(\mu, \sigma^2)$ where $\bar{Y}^* \sim \text{Lap}(\bar{Y}, 1/(\varepsilon_1 n))$ and $S^{2*} \sim \text{Lap}(S^2, 1/(\varepsilon_2 n))$ are released, if an analyst has prior $p(\mu, \sigma^2) \propto (\sigma^2)^{-1}$, then $p(\mu, \sigma^2 \mid \bar{Y}^*, S^{2*})$ is not a proper probability distribution.

Proof. To begin, note that under the prior $p(\mu, \sigma^2) \propto (\sigma^2)^{-1}$, for observed $\bar{Y}^* = \bar{y}^*$ and $S^{2*} = s^{2*}$ the posterior distribution, if proper, should be

$$p(\mu, \sigma^2 \mid \bar{y}^*, s^{2*}) = \frac{p(\bar{y}^*, s^{2*} \mid \mu, \sigma^2) p(\mu, \sigma^2)}{\int_0^\infty \int_{-\infty}^\infty p(\bar{y}^*, s^{2*} \mid \mu, \sigma^2) p(\mu, \sigma^2) d\mu d\sigma^2} \quad (\text{F.1})$$

$$= \frac{p(\bar{y}^*, s^{2*} \mid \mu, \sigma^2) (\sigma^2)^{-1}}{\int_0^\infty \int_{-\infty}^\infty p(\bar{y}^*, s^{2*} \mid \mu, \sigma^2) (\sigma^2)^{-1} d\mu d\sigma^2}. \quad (\text{F.2})$$

A proper probability distribution requires the integral in the denominator of (F.2) to be finite. We now examine this integral. Introducing the latent variables $\bar{Y} = \bar{y}$ and $S^2 = s^2$

and exploiting the conditional independence structure, this integral is equivalent to

$$\int_0^\infty \int_{-\infty}^\infty p(\bar{y}^*, s^{2*} \mid \mu, \sigma^2) (\sigma^2)^{-1} d\mu d\sigma^2 \quad (\text{F.3})$$

$$= \int_0^\infty \int_{-\infty}^\infty \int_0^\infty \int_{-\infty}^\infty p(\bar{y}^*, s^{2*}, \bar{y}, s^2 \mid \mu, \sigma^2) (\sigma^2)^{-1} d\bar{y} ds^2 d\mu d\sigma^2 \quad (\text{F.4})$$

$$= \int_0^\infty (\sigma^2)^{-1} \int_{-\infty}^\infty \int_0^\infty \int_{-\infty}^\infty p(\bar{y}^* \mid \bar{y}) p(s^{2*} \mid s^2) p(\bar{y} \mid \mu, \sigma^2) p(s^2 \mid \sigma^2) d\bar{y} ds^2 d\mu d\sigma^2 \quad (\text{F.5})$$

$$= \int_0^\infty (\sigma^2)^{-1} \left[\int_0^\infty p(s^{2*} \mid s^2) p(s^2 \mid \sigma^2) ds^2 \right] \left[\int_{-\infty}^\infty \int_{-\infty}^\infty p(\bar{y}^* \mid \bar{y}) p(\bar{y} \mid \mu, \sigma^2) d\bar{y} d\mu \right] d\sigma^2. \quad (\text{F.6})$$

We examine the double integral in the second bracket. If we exchange the order of integration,

$$\int_{-\infty}^\infty \int_{-\infty}^\infty p(\bar{y}^* \mid \bar{y}) p(\bar{y} \mid \mu, \sigma^2) d\mu d\bar{y} = \int_{-\infty}^\infty p(\bar{y}^* \mid \bar{y}) \int_{-\infty}^\infty \frac{1}{\sqrt{2\pi\sigma^2}} e^{-\frac{(\bar{y}-\mu)^2}{2\sigma^2}} d\mu d\bar{y}, \quad (\text{F.7})$$

since we recognize the quantity under the inner integral as the density function of $\mathcal{N}(\bar{y}, \sigma^2)$, the integral is equal to 1. Thus,

$$\int_{-\infty}^\infty \int_{-\infty}^\infty p(\bar{y}^* \mid \bar{y}) p(\bar{y} \mid \mu, \sigma^2) d\mu d\bar{y} = \int_{-\infty}^\infty p(\bar{y}^* \mid \bar{y}) d\bar{y} = \int_{-\infty}^\infty \frac{\varepsilon_1 n}{2} e^{-\varepsilon_1 n |\bar{y}^* - \bar{y}|} d\bar{y} = 1, \quad (\text{F.8})$$

since we recognize the quantity under the integral as the density function of $\text{Lap}(\bar{y}^*, 1/(\varepsilon_1 n))$. By Fubini's Theorem, since this integral is finite and the quantity being integrated is nonnegative, it is equivalent to the quantity in brackets above. Thus, the quantity of interest reduces to

$$\int_0^\infty \int_{-\infty}^\infty p(\bar{y}^*, s^{2*} \mid \mu, \sigma^2) (\sigma^2)^{-1} d\mu d\sigma^2 = \int_0^\infty (\sigma^2)^{-1} \int_0^\infty p(s^{2*} \mid s^2) p(s^2 \mid \sigma^2) ds^2 d\sigma^2. \quad (\text{F.9})$$

If $s^{2*} > 0$, then we bound the inner integral below as follows.

$$\int_0^\infty p(s^{2*} | s^2) p(s^2 | \sigma^2) ds^2 = \frac{\varepsilon_2 n \left(\frac{n-1}{2\sigma^2}\right)^{\frac{n-1}{2}}}{2\Gamma(\frac{n-1}{2})} \int_0^\infty e^{-\varepsilon_2 n |s^{2*} - s^2|} (s^2)^{\frac{n-1}{2}-1} e^{-\frac{n-1}{2\sigma^2} s^2} ds^2 \quad (\text{F.10})$$

$$\begin{aligned} &= \frac{\varepsilon_2 n \left(\frac{n-1}{2\sigma^2}\right)^{\frac{n-1}{2}}}{2\Gamma(\frac{n-1}{2})} \left[e^{-\varepsilon_2 n s^{2*}} \int_0^{s^{2*}} (s^2)^{\frac{n-1}{2}-1} e^{-\left(\frac{n-1}{2\sigma^2} - \varepsilon_2 n\right) s^2} ds^2 \right. \\ &\quad \left. + e^{\varepsilon_2 n s^{2*}} \int_{s^{2*}}^\infty (s^2)^{\frac{n-1}{2}-1} e^{-\left(\frac{n-1}{2\sigma^2} + \varepsilon_2 n\right) s^2} ds^2 \right] \quad (\text{F.11}) \end{aligned}$$

$$\begin{aligned} &\geq \frac{\varepsilon_2 n \left(\frac{n-1}{2\sigma^2}\right)^{\frac{n-1}{2}}}{2\Gamma(\frac{n-1}{2})} \left[e^{-\varepsilon_2 n s^{2*}} \int_0^{s^{2*}} (s^2)^{\frac{n-1}{2}-1} e^{-\left(\frac{n-1}{2\sigma^2} + \varepsilon_2 n\right) s^2} ds^2 \right. \\ &\quad \left. + e^{-\varepsilon_2 n s^{2*}} \int_{s^{2*}}^\infty (s^2)^{\frac{n-1}{2}-1} e^{-\left(\frac{n-1}{2\sigma^2} + \varepsilon_2 n\right) s^2} ds^2 \right] \quad (\text{F.12}) \end{aligned}$$

$$= \frac{\varepsilon_2 n \left(\frac{n-1}{2\sigma^2}\right)^{\frac{n-1}{2}}}{2\Gamma(\frac{n-1}{2})} e^{-\varepsilon_2 n s^{2*}} \int_0^\infty (s^2)^{\frac{n-1}{2}-1} e^{-\left(\frac{n-1}{2\sigma^2} + \varepsilon_2 n\right) s^2} ds^2 \quad (\text{F.13})$$

$$= \frac{\varepsilon_2 n}{2} \left(\frac{\frac{n-1}{2\sigma^2}}{\frac{n-1}{2\sigma^2} + \varepsilon_2 n} \right)^{\frac{n-1}{2}} \quad (\text{F.14})$$

$$= \frac{\varepsilon_2 n}{2} \left(\frac{\frac{n-1}{2\varepsilon_2 n}}{\frac{n-1}{2\varepsilon_2 n} + \sigma^2} \right)^{\frac{n-1}{2}}. \quad (\text{F.15})$$

The inequality follows since $e^x \geq e^{-x}$ for $x > 0$ and since $e^{-(a-b)x} \geq e^{-(a+b)x}$ for $a > b > 0$ and $x > 0$. If $s^{2*} \leq 0$, then the same bound holds by an analogous argument (with the integral from 0 to s^{2*} omitted). Thus, the quantity of interest is bounded as follows.

$$\int_0^\infty \int_{-\infty}^\infty p(\bar{y}^*, s^{2*} | \mu, \sigma^2) (\sigma^2)^{-1} d\mu d\sigma^2 \geq \frac{\varepsilon_2 n}{2} \int_0^\infty (\sigma^2)^{-1} \left(\frac{\frac{n-1}{2\varepsilon_2 n}}{\frac{n-1}{2\varepsilon_2 n} + \sigma^2} \right)^{\frac{n-1}{2}} d\sigma^2 \quad (\text{F.16})$$

$$\geq \frac{\varepsilon_2 n}{2} \int_0^1 (\sigma^2)^{-1} \left(\frac{\frac{n-1}{2\varepsilon_2 n}}{\frac{n-1}{2\varepsilon_2 n} + \sigma^2} \right)^{\frac{n-1}{2}} d\sigma^2 \quad (\text{F.17})$$

$$\geq \frac{\varepsilon_2 n}{2} \int_0^1 (\sigma^2)^{-1} \left(\frac{\frac{n-1}{2\varepsilon_2 n}}{\frac{n-1}{2\varepsilon_2 n} + 1} \right)^{\frac{n-1}{2}} d\sigma^2 \quad (\text{F.18})$$

$$= \frac{\varepsilon_2 n}{2} \left(\frac{\frac{n-1}{2\varepsilon_2 n}}{\frac{n-1}{2\varepsilon_2 n} + 1} \right)^{\frac{n-1}{2}} \int_0^1 (\sigma^2)^{-1} d\sigma^2. \quad (\text{F.19})$$

Since this integral diverges, the quantity of interest diverges. Thus, (F.2) is not a proper probability distribution. $\square$

Theorem 4.2. *For confidential data $Y_1, \dots, Y_n \stackrel{iid}{\sim} \mathcal{N}(\mu, \sigma^2)$ where $n > 3$ and $\bar{Y}^* \sim \text{Lap}(\bar{Y}, 1/(\varepsilon_1 n))$ and $S^{2*} \sim \text{Lap}(S^2, 1/(\varepsilon_2 n))$ are released, if an analyst has prior $p(\mu, \sigma^2) \propto 1$, then $p(\mu, \sigma^2 \mid \bar{Y}^*, S^{2*})$ is a proper probability distribution.*

Proof. To begin, note that under the prior $p(\mu, \sigma^2) \propto 1$, for observed $\bar{Y}^* = \bar{y}^*$ and $S^{2*} = s^{2*}$ the posterior distribution, if proper, should be

$$p(\mu, \sigma^2 \mid \bar{y}^*, s^{2*}) = \frac{p(\bar{y}^*, s^{2*} \mid \mu, \sigma^2) p(\mu, \sigma^2)}{\int_0^\infty \int_{-\infty}^\infty p(\bar{y}^*, s^{2*} \mid \mu, \sigma^2) p(\mu, \sigma^2) d\mu d\sigma^2} \quad (\text{F.20})$$

$$= \frac{p(\bar{y}^*, s^{2*} \mid \mu, \sigma^2)}{\int_0^\infty \int_{-\infty}^\infty p(\bar{y}^*, s^{2*} \mid \mu, \sigma^2) d\mu d\sigma^2}. \quad (\text{F.21})$$

The posterior is a proper probability distribution if the integral in the denominator of (F.21) is finite. By analogy to the proof of Theorem 4.1, this integral is equivalent to

$$\begin{aligned} & \int_0^\infty \int_{-\infty}^\infty p(\bar{y}^*, s^{2*} \mid \mu, \sigma^2) d\mu d\sigma^2 \\ &= \int_0^\infty \left[\int_0^\infty p(s^{2*} \mid s^2) p(s^2 \mid \sigma^2) ds^2 \right] \left[\int_{-\infty}^\infty \int_{-\infty}^\infty p(\bar{y}^* \mid \bar{y}) p(\bar{y} \mid \mu, \sigma^2) d\bar{y} d\mu \right] d\sigma^2. \end{aligned} \quad (\text{F.22})$$

$$(\text{F.23})$$

By (F.8), the double integral in the second bracket is equal to 1. Thus, the quantity of interest reduces to

$$\int_0^\infty \int_{-\infty}^\infty p(\bar{y}^*, s^{2*} \mid \mu, \sigma^2) d\mu d\sigma^2 = \int_0^\infty \int_0^\infty p(s^{2*} \mid s^2) p(s^2 \mid \sigma^2) ds^2 d\sigma^2. \quad (\text{F.24})$$

Exchanging the order of integration yields

$$\int_0^\infty p(s^{2*} \mid s^2) \int_0^\infty p(s^2 \mid \sigma^2) d\sigma^2 ds^2. \quad (\text{F.25})$$

Since $n > 3$, the inner integral is the kernel of an inverse-gamma distribution and so

$$\int_0^\infty p(s^2 \mid \sigma^2) d\sigma^2 = \int_0^\infty \frac{\left(\frac{n-1}{2}\right)^{\frac{n-1}{2}}}{\Gamma(\frac{n-1}{2})} (s^2)^{\frac{n-1}{2}-1} e^{-\frac{(n-1)s^2}{2\sigma^2}} d\sigma^2 \quad (\text{F.26})$$

$$= \frac{\left(\frac{n-1}{2}\right)^{\frac{n-1}{2}}}{\Gamma(\frac{n-1}{2})} (s^2)^{\frac{n-1}{2}-1} \int_0^\infty (\sigma^2)^{-\frac{n-3}{2}-1} e^{-\frac{(n-1)s^2}{2\sigma^2}} d\sigma^2 \quad (\text{F.27})$$

$$= \frac{\left(\frac{n-1}{2}\right)^{\frac{n-1}{2}}}{\Gamma(\frac{n-1}{2})} (s^2)^{\frac{n-1}{2}-1} \frac{\Gamma(\frac{n-3}{2})}{\left(\frac{(n-1)s^2}{2}\right)^{\frac{n-3}{2}}} \quad (\text{F.28})$$

$$= \frac{\frac{n-1}{2} \Gamma(\frac{n-3}{2})}{\Gamma(\frac{n-1}{2})} \quad (\text{F.29})$$

$$= \frac{n-1}{n-3}. \quad (\text{F.30})$$

Thus,

$$\int_0^\infty p(s^{2*} | s^2) \int_0^\infty p(s^2 | \sigma^2) d\sigma^2 ds^2 = \frac{n-1}{n-3} \int_0^\infty \frac{\varepsilon_2 n}{2} e^{-\varepsilon_2 n |s^{2*} - s^2|} ds^2 \quad (\text{F.31})$$

$$= \frac{n-1}{n-3} \left(\frac{1}{2} + \frac{1}{2} \text{sgn}(s^{2*}) (1 - e^{-\varepsilon_2 n |s^{2*}|}) \right), \quad (\text{F.32})$$

since we recognize the quantity under the integral as the density function of $\text{Lap}(s^{2*}, 1/(\varepsilon_2 n))$, we plug in the form of its CDF at zero. By Fubini's Theorem, since this integral is finite and the quantity being integrated is nonnegative, it is equivalent to the original quantity of interest. Thus,

$$\int_0^\infty \int_{-\infty}^\infty p(\bar{y}^*, s^{2*} | \mu, \sigma^2) (\sigma^2)^0 d\mu d\sigma^2 = \frac{n-1}{n-3} \left(\frac{1}{2} + \frac{1}{2} \text{sgn}(s^{2*}) (1 - e^{-\varepsilon_2 n |s^{2*}|}) \right) < \infty \quad (\text{F.33})$$

and so the posterior is a proper probability distribution. $\square$

APPENDIX G. NONPRIVATE LAPLACE-UNIFORM POSTERIOR DISTRIBUTION

In this section, we demonstrate that if $X \sim \text{Lap}(\xi, 1/\lambda)$ and λ is known, then a uniform prior on ξ yields a proper, frequentist matching posterior. The main result is as follows.

Theorem G.1. *Let $X \sim \text{Lap}(\xi, 1/\lambda)$, where λ is known. Then the prior $p(\xi) \propto 1$ yields a proper posterior distribution that is frequentist matching.*

Proof. Applying properties of the Laplace distribution, if $X \sim \text{Lap}(\xi, 1/\lambda)$, then $\xi - X \sim \text{Lap}(0, 1/\lambda)$. Note that this is a pivotal quantity, since its distribution does not depend on any unknown parameters. Let $\ell_{\alpha/2}$ and $\ell_{1-\alpha/2}$ be the $\alpha/2$ and $(1 - \alpha/2)$ quantiles of $\text{Lap}(\xi, 1/\lambda)$. Then,

$$P[X + \ell_{\alpha/2} \leq \xi \leq X + \ell_{1-\alpha/2}] = P[\ell_{\alpha/2} \leq \xi - X \leq \ell_{1-\alpha/2}] = 1 - \alpha \quad (\text{G.1})$$

and so $[X + \ell_{\alpha/2}, X + \ell_{1-\alpha/2}]$ is a $(1 - \alpha)$ confidence interval for ξ .

In the Bayesian setting, if $p(\xi) \propto 1$ then

$$p(\xi | X) \propto p(X | \xi) p(\xi) \propto \exp\{-\lambda|X - \xi|\} = \exp\{-\lambda|\xi - X|\}.$$

Thus, the posterior distribution is $(\xi | X) \sim \text{Lap}(X, 1/\lambda)$. Again applying properties of the Laplace distribution, $(\xi - X | X) \sim \text{Lap}(0, 1/\lambda)$. Thus,

$$P[X + \ell_{\alpha/2} \leq \xi \leq X + \ell_{1-\alpha/2} | X] = P[\ell_{\alpha/2} \leq \xi - X \leq \ell_{1-\alpha/2} | X] = 1 - \alpha \quad (\text{G.2})$$

and so $[X + \ell_{\alpha/2}, X + \ell_{1-\alpha/2}]$ is a $(1 - \alpha)$ credible interval. Since this interval is the same as the confidence interval above, this posterior distribution is frequentist matching. $\square$

APPENDIX H. CONSTRAINTS IN SIMPLE LINEAR REGRESSION

In this section, we derive constraints on the model proposed by [Bernstein and Sheldon \(2019\)](#) in the simple linear regression setting. Letting $\mathbf{X} = [\mathbf{1}_n \ \mathbf{x}_1]$, the authors release $\mathbf{X}^\top \mathbf{X}$, $\mathbf{X}^\top \mathbf{Y}$, and $\mathbf{Y}^\top \mathbf{Y}$, which include the following elements.

$$\mathbf{X}^\top \mathbf{X} = \begin{bmatrix} \mathbf{1}_n^\top \mathbf{1}_n & \mathbf{1}_n^\top \mathbf{x}_1 \\ \mathbf{x}_1^\top \mathbf{1}_n & \mathbf{x}_1^\top \mathbf{x}_1 \end{bmatrix}, \quad \mathbf{X}^\top \mathbf{Y} = \begin{bmatrix} \mathbf{1}_n^\top \mathbf{Y} \\ \mathbf{x}_1^\top \mathbf{Y} \end{bmatrix}. \quad (\text{H.1})$$

Note that $\mathbf{1}_n^\top \mathbf{1}_n = n$ and as sums of squares, $\mathbf{x}_1^\top \mathbf{x}_1 = \sum_{i=1}^n x_i^2 \geq 0$ and $\mathbf{Y}^\top \mathbf{Y} = \sum_{i=1}^n y_i^2 \geq 0$. The requirement that $x_i, y_i \in [0, 1]$ for all $i \in \{1, \dots, n\}$ implies the following additional constraints.

- (1) Since $x_i \leq 1$ for all i , it follows that $\mathbf{x}_1^\top \mathbf{1}_n = \mathbf{1}_n^\top \mathbf{x}_1 = \sum_{i=1}^n x_i \leq n$.
- (2) Since $y_i \leq 1$ for all i , it follows that $\mathbf{1}_n^\top \mathbf{Y} = \sum_{i=1}^n y_i \leq n$.
- (3) Since $x_i, y_i \geq 0$ for all i , it follows that $\mathbf{x}_1^\top \mathbf{Y} = \sum_{i=1}^n x_i y_i \geq 0$.
- (4) Since $x_i \leq 1$ for all i , it follows that $\mathbf{x}_1^\top \mathbf{Y} = \sum_{i=1}^n x_i y_i \leq \sum_{i=1}^n y_i = \mathbf{1}_n^\top \mathbf{Y}$.
- (5) Since $y_i \leq 1$ for all i , it follows that $\mathbf{x}_1^\top \mathbf{Y} = \sum_{i=1}^n x_i y_i \leq \sum_{i=1}^n x_i = \mathbf{1}_n^\top \mathbf{x}_1$.
- (6) Since $0 \leq x_i \leq 1$ for all i , $x_i^2 \leq x_i$ and so $\mathbf{x}_1^\top \mathbf{x}_1 = \sum_{i=1}^n x_i^2 \leq \sum_{i=1}^n x_i = \mathbf{1}_n^\top \mathbf{x}_1$.
- (7) Since $0 \leq y_i \leq 1$ for all i , $y_i^2 \leq y_i$ and so $\mathbf{Y}^\top \mathbf{Y} = \sum_{i=1}^n y_i^2 \leq \sum_{i=1}^n y_i = \mathbf{1}_n^\top \mathbf{Y}$.

The form of the model implies constraints on the parameters θ_0 , θ_1 , and σ^2 . We can rewrite the model in the following form, where for each $i \in \{1, \dots, n\}$,

$$y_i \sim \mathcal{N}(\theta_0 + \theta_1 x_i, \sigma^2). \quad (\text{H.2})$$

For the regression coefficients, θ_0 and θ_1 , we observe that by monotonicity of expectation, since $y_i \in [0, 1]$, it follows that $E[y_i] = \theta_0 + \theta_1 x_i \in [0, 1]$. Using this fact, for a given value of θ_1 , we obtain the following constraints on θ_0 .

- (1) When $\theta_1 \geq 0$, since $x_i \geq 0$ for all i , it follows that $\theta_0 \leq \theta_0 + \theta_1 x_i \leq 1$.
- (2) When $\theta_1 \geq 0$, since $x_i \leq 1$ for all i , it follows that $\theta_0 + \theta_1 \geq \theta_0 + \theta_1 x_i \geq 0$ and so $\theta_0 \geq -\theta_1$.
- (3) When $\theta_1 < 0$, since $x_i \geq 0$ for all i , it follows that $\theta_0 \geq \theta_0 + \theta_1 x_i \geq 0$.
- (4) When $\theta_1 < 0$, since $x_i \leq 1$ for all i , it follows that $\theta_0 + \theta_1 \leq \theta_0 + \theta_1 x_i \leq 1$ and so $\theta_0 \leq 1 - \theta_1$.

We summarize these constraints as follows.

$$\begin{cases} 0 \leq \theta_0 \leq 1 - \theta_1, & \text{if } \theta_1 < 0; \\ -\theta_1 \leq \theta_0 \leq 1, & \text{if } \theta_1 \geq 0. \end{cases} \quad (\text{H.3})$$

Finally, we use a similar argument to [Observation 3.4](#) to obtain a constraint on σ^2 . By monotonicity of expectation and since $y_i \in [0, 1]$ implies $y_i^2 \leq y_i$, for all i ,

$$\sigma^2 = \text{Var}[y_i] = E[y_i^2] - E[y_i]^2 \leq E[y_i] - E[y_i]^2 \leq 1/4. \quad (\text{H.4})$$

The final inequality is obtained by observing that the quadratic $f(x) = x - x^2$ is minimized at $x = 1/2$ and has minimum $f(1/2) = 1/4$. Since the variance of y_i must be nonnegative, it follows that $\sigma^2 \in [0, 1/4]$.

APPENDIX I. PRIOR DISTRIBUTIONS FOR GAUSSIAN MODELS WITH DIFFERENTIAL PRIVACY

This appendix illustrates an approach to incorporating bounds on the data values when fitting a Gaussian model via the Gibbs sampling framework developed by Ju et al. (2022). Their sampler represents a general strategy for Bayesian inference. As the authors state, it can “overcome the intractable marginal likelihood resulting from privatization, and is applicable to a wide range of statistical models and privacy mechanisms” (Ju et al., 2022). Their sampler is applicable whenever the analyst can sample plausible values from the posited model for the confidential data and can compute the density associated with the privacy mechanism (Ju et al., 2022). These conditions are met in our setting.

As discussed in the main body, we presume that the analyst considers the univariate Gaussian distribution as the model they would use on the confidential data absent privacy requirements, e.g., because the truncation limits are far enough in the tails as to be practically irrelevant for inferences with the confidential data (which is needed to make a Gaussian model defensible in the first place). This analyst interprets μ and σ^2 as the expectation and variance of the data generating process.

As a point of clarity, our illustration does not consider an analyst who presumes the underlying confidential data follow a truncated Gaussian distribution. In the truncated Gaussian distribution, the parameters (μ, σ^2) do not necessarily correspond to the expectation and variance of the data generating distribution. Thus, (μ, σ^2) need not necessarily fall in the feasible region A ; in fact, the marginal prior and posterior distributions can have support over all real numbers for μ and over all positive real numbers for σ^2 . Analysts who presume truncated Gaussian models would need to transform the posterior distribution for (μ, σ^2) into a posterior distribution for the expectation and variance using the properties of the truncated Gaussian distribution (Johnson et al., 1994, Section 10.1). Our methodology for enforcing constraints on the support of the prior distributions is not well-suited for the truncated Gaussian model, as there is no theoretical justification for incorporating constraints on (μ, σ^2) in this model.

I.1. Constraints on Gaussian Model Parameters Using an Adaptation of the Sampler of Ju et al. (2022). As discussed in the main body, an analyst who seeks to estimate the parameters of a Gaussian model may wish to enforce the constraints on the unobserved $(\bar{Y}, S^2)$ and (μ, σ^2) that arise from bounds on the data values. We now illustrate how to incorporate such bounds using a modified version of the sampler in Ju et al. (2022). In doing so, we constrain the prior support of (μ, σ^2) to the feasible region A , for example, via the default prior distribution of Section 4.

I.1.1. The Sampler. To provide context, we first describe the Gibbs sampler proposed by Ju et al. (2022) for the univariate Gaussian setting where we do not incorporate constraints on data values or model parameters. We then outline how to adapt this sampler to incorporate constraints on the parameter values and data. Throughout, we presume that the analyst would like to specify a Gaussian model (not a truncated Gaussian model) for the confidential data generating distribution.

Review of Sampler Without Constraints. Disregarding all constraints, we can use the sampler of Ju et al. (2022) to estimate the posterior distribution of the univariate Gaussian model with the prior distribution in Equation (3.1). At the t^{th} iteration, we sample from $n + 2$ full conditionals: those of μ , σ^2 , and each of the confidential Y_i . The full conditionals for μ and σ^2 are identical to the ones used in the public setting; see Section I.1.2. For each Y_i , a proposal is drawn from $Y_i' \sim \mathcal{N}(\mu_{(t)}, \sigma_{(t)}^2)$. Let $\bar{Y}_{\text{prev}}$ and $\bar{Y}'$ be the sample mean with and without the proposal, respectively, i.e.,

$$\bar{Y}_{\text{prev}} = \frac{1}{n} \left(\sum_{j=1}^{i-1} Y_{i(t)} + Y_{i(t-1)} + \sum_{j=i+1}^n Y_{i(t-1)} \right), \quad \bar{Y}' = \frac{1}{n} \left(\sum_{j=1}^{i-1} Y_{i(t)} + Y_i' + \sum_{j=i+1}^n Y_{i(t-1)} \right). \quad (\text{I.1})$$

Let S_{prev}^2 and $S^{2'}$ be defined similarly for the sample variance. Notably, for a given proposal Y_i' , the updated moments $(\bar{Y}', S^{2'})$ can be computed from $(\bar{Y}_{\text{prev}}, S_{\text{prev}}^2)$ in constant time. Then, the proposed Y_i is accepted with probability,

$$r = \min \left\{ 1, \exp \left[-\varepsilon_1 n (|\bar{Y}^* - \bar{Y}'| - |\bar{Y}^* - \bar{Y}_{\text{prev}}|) - \varepsilon_2 n (|S^{2*} - S^{2'}| - |S^{2*} - S_{\text{prev}}^2|) \right] \right\}. \quad (\text{I.2})$$

One can verify empirically that the estimated posterior distribution from this Gibbs sampler is identical to the posterior distribution from Section 3.1 (using the exact full conditional for S^2) without constraining the sample moments or parameters. Thus, as noted in Section 3.2 and Section 4, we potentially can improve accuracy by incorporating the constraints on sufficient statistics and parameters implied by the data bounds.

Accounting for Constraints. We now summarize at a high level our modifications of the sampler in Appendix to enforce the constraint that each $Y_i \in [a, b]$. The basic idea is to apply methods similar to those described in Section 3.2. For μ and σ^2 , we constrain the full conditionals to lie in A . We also specify a proper prior distribution, which is a required assumption to utilize the sampler of Ju et al. (2022). For the illustrations below, we utilize the proper uniform prior over the set A defined in Equation (3.4). For each potential draw of Y_i , we enforce the bounds as part of the sampler by rejecting any proposed draws outside of $[a, b]$. Details of the Gibbs sampler and the forms of the full conditionals are provided in Section I.1.2.

I.1.2. Modified Sampler Incorporating Constraints. The modified sampler, which we abbreviate as MJS, is presented in Algorithm 3. The first step is to initialize values for each of the variables, for which the considerations are similar to those of Algorithm 2. The analyst then alternates between sampling from the full conditionals for μ, σ^2 and using a Metropolis-Hastings update for the Y_i . Similarly to the measurement error model, the full conditionals for μ and σ^2 take the forms from the public setting, truncated to the appropriate region. That is, for a constrained, informative prior of the form in Equation (3.5), the full conditionals are

$$\mu \mid \sigma^2, Y_1, \dots, Y_n, \bar{Y}^*, S^{2*} \sim \mathcal{N}_{[1/2 - \sqrt{1/4 - \sigma^2}, 1/2 + \sqrt{1/4 - \sigma^2}]} \left(\frac{n\bar{Y} + \kappa_0\mu_0}{n + \kappa_0}, \frac{\sigma^2}{n + \kappa_0} \right) \quad (\text{I.3})$$

$$\sigma^2 \mid \mu, Y_1, \dots, Y_n, \bar{Y}^*, S^{2*} \sim \text{IG}_{[0, \mu(1-\mu)]} \left(\frac{n + \nu_0}{2}, \frac{\nu_0\sigma_0^2 + (n-1)S^2 + n(\bar{Y} - \mu)^2}{2} \right). \quad (\text{I.4})$$

Algorithm 3: Run the MJS sampler

Input: Data $(\bar{Y}^*, S^{2*})$, privacy budget $(\varepsilon_1, \varepsilon_2)$, sample size n , iterations T , hyperparameters $(\mu_0, \sigma_0^2, \kappa_0, \nu_0)$
Initialize $\sigma_{(0)}^2, Y_{1(0)}, \dots, Y_{n(0)}$
for t **in** $1 : T$ **do**
 Sample $\mu_{(t)}$ from $(\mu \mid \sigma^2 = \sigma_{(t-1)}^2, Y_1 = Y_{1(t-1)}, \dots, Y_n = Y_{n(t-1)}, \bar{Y}^*, S^{2*})$
 Sample $\sigma_{(t)}^2$ from $(\sigma^2 \mid \mu = \mu_{(t)}, Y_1 = Y_{1(t-1)}, \dots, Y_n = Y_{n(t-1)}, \bar{Y}^*, S^{2*})$
 Sample $Y'_1, \dots, Y'_n \stackrel{iid}{\sim} \mathcal{N}(\mu_{(t)}, \sigma_{(t)}^2)$
 for i **in** $1 : n$ **do**
 if $0 \leq Y'_i \leq 1$ **then**
 Compute $\bar{Y}_{\text{prev}}, S_{\text{prev}}^2$ and $\bar{Y}', S^{2'}$ as described in (I.1)
 Set $r = \min \left\{ 1, e^{-\varepsilon_1 n (|\bar{Y}^* - \bar{Y}'| - |\bar{Y}^* - \bar{Y}_{\text{prev}}|) - \varepsilon_2 n (|S^{2*} - S^{2'}| - |S^{2*} - S_{\text{prev}}^2|)} \right\}$
 else
 Set $r = 0$
 Set $Y_{i(t)} = \begin{cases} Y'_i, & \text{with probability } r; \\ Y_{i(t-1)}, & \text{with probability } 1 - r. \end{cases}$
Output: $(\mu_{(1)}, \dots, \mu_{(T)}), (\sigma_{(1)}^2, \dots, \sigma_{(T)}^2)$

For the prior that is uniform over the set A defined in Equation (3.4), we have

$$\mu \mid \sigma^2, \omega^2, \bar{Y}, S^2, \bar{Y}^*, S^{2*} \sim \mathcal{N}_{[1/2 - \sqrt{1/4 - \sigma^2}, 1/2 + \sqrt{1/4 - \sigma^2}]} \left(\bar{Y}, \frac{\sigma^2}{n} \right) \quad (\text{I.5})$$

$$\sigma^2 \mid \mu, \omega^2, \bar{Y}, S^2, \bar{Y}^*, S^{2*} \sim \text{IG}_{[0, \mu(1-\mu)]} \left(\frac{n-2}{2}, \frac{(n-1)S^2 + n(\bar{Y} - \mu)^2}{2} \right). \quad (\text{I.6})$$

See Appendix C.3 for further details on deriving full conditionals with the constraint that $(\mu, \sigma^2) \in A$.

For the Metropolis-Hastings update, any sample outside the range $[0, 1]$ for each Y_i is rejected. Otherwise, the sample is accepted with probability r as defined in (I.2). The final step is to output the posterior draws for the two parameters.

Sampling Y_i subject to bounds is done to ensure the draws of the sufficient statistics in the MCMC chains satisfy the constraints. However, as noted by a reviewer, this convenience creates a type of incoherency in the full conditionals. Specifically, the full conditional for each Y_i follows a truncated distribution. However, the forms of the full conditionals of (μ, σ^2) derive from a Gaussian distribution. One could resolve this incoherency by in fact switching the entire model to a truncated Gaussian model. In this case, we are no longer in the setting where an analyst seeks to estimate a Gaussian model with μ interpreted as the mean and σ^2 interpreted as the variance of the observations. Thus, despite the potential incoherency introduced by the computational convenience to enforce the constraints, we evaluate this sampling algorithm as an approximation for estimating the Gaussian model without truncation using a different approach than the measurement error framework of the main body.

To run this sampler without constraints, as described in Section I.1.1, Algorithm 3 changes in two ways. First, the requirement that $Y'_i \in [0, 1]$ is omitted within the inner

for loop. Second, the truncation is removed from the full conditional distributions for μ and σ^2 in (I.3)–(I.6); samples are drawn from the traditional Gaussian or inverse-gamma distribution.

I.1.3. Evaluating the Constrained Posterior Distributions. Here we evaluate the inferential performance of the MJS sampler as a way to incorporate constraints on the data values introduced by DP. Alongside the MJS sampler, we consider the performance of the constrained sufficient statistics sampler from Section 3.2.1, which we refer to as CSS. We emphasize that our purpose is not to find the “best” approach for accounting for constraints in Gaussian models; rather, it is to emphasize the benefits of incorporating constraints and present methods to do so.

Consider a setting with sample size $n = 50$ and privacy budget $\varepsilon_1 = \varepsilon_2 = 0.25$, where we enforce $Y_i \in [0, 1]$ for all i to implement the Laplace mechanism. As a first illustration, we suppose that $\bar{Y} = 0.5$ is in the center of the constrained region and $S^2 = 0.2^2$ is fairly large. Figure 11 plots the posterior distribution of $(\mu, \sigma^2 \mid \bar{Y}^*, S^{2*})$ for a particular draw of $(\bar{Y}^*, S^{2*})$ from the Laplace mechanism. As shown in the left panel of Figure 11, CSS yields a posterior distribution with mode around the released values. The posterior has considerable uncertainty and places probability density throughout the feasible region. As shown in the right panel of Figure 11, MJS also yields a posterior with its estimate for μ close to the released $\bar{Y}^*$. Its estimate for σ^2 is smaller than the released S^{2*} . The posterior distribution has less uncertainty, concentrating the probability density in the region where σ^2 is small and μ is near $\bar{Y}^*$.

The difference in the two posterior distributions stems from the implications of constructing draws of $(\bar{Y}, S^2)$ by individually sampling constrained Y_i , as in MJS, or sampling directly from the sample moments’ full conditional distributions, as in CSS. As shown in Observation 3.5, any data bounded in $[a, b]$ has bounded sample variance, S^2 . Since MJS favors values of Y_i that have high density under the Gaussian distribution, i.e., are far from the bounds a and b , the Y_i have a sample variance in the lower part of the feasible range for S^2 . As a result, MJS favors relatively small values of σ^2 , as observed in Figure 11. In contrast, CSS gives posterior density to values of S^2 in the upper part of the feasible range A , as it is not tied to finding Gaussian models that allow for high density samples of each individual Y_i .

The tendency of MJS to favor values of (μ, σ^2) that generate samples of Y_i relatively far from $[a, b]$ can have advantages in situations where in fact the data values are far from the bounds. However, this tendency can be problematic when data are close to the bounds, as we now illustrate. Consider the same dataset with the same privacy budget, except suppose the data curator uses much wider bounds to implement the Laplace mechanism: all $Y_i \in [0, 25]$. We keep the observed statistics as $\bar{Y} = 0.5$ and $S^2 = 0.2^2$, which are now very close to the boundary of the feasible region.

Figure 12 displays the posterior distributions for the same draw from the Laplace mechanism as in Figure 11 (with updated scale parameter). CSS yields a posterior distribution with its mode near the boundary. Since $\bar{Y}^*$ is negative, it focuses the probability density in the part of the feasible region closest to the released value, with substantial uncertainty. MJS still yields a posterior distribution with substantial probability density near the center of the feasible region. The posterior distribution places fairly little probability density in the region where μ is near zero, despite this being the most likely region for $\bar{Y}$ under the Laplace mechanism when the released $\bar{Y}^*$ is negative. As discussed above, the need to sample Y_i

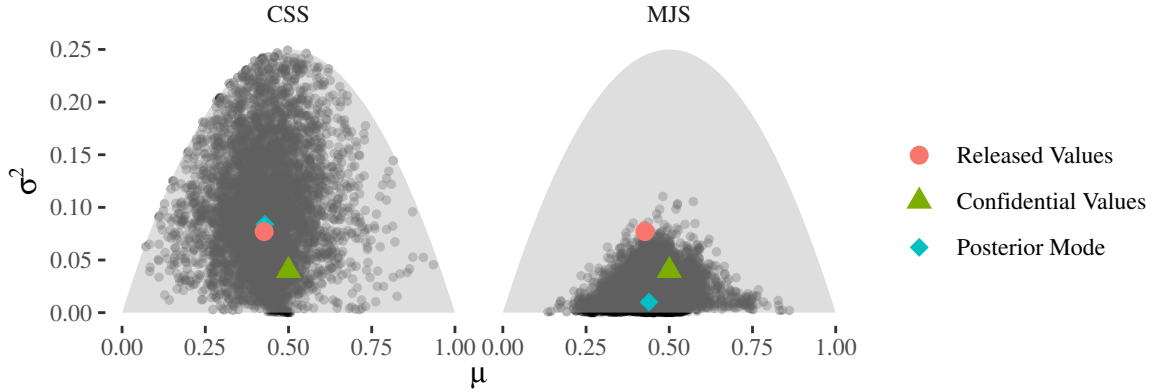


Figure 11: Posterior distribution for (μ, σ^2) for CSS (left panel) and MJS (right panel), both using the constrained uniform prior distribution. Results use $n = 50$, $\varepsilon_1 = \varepsilon_2 = 0.25$, $[a, b] = [0, 1]$. The unreleased confidential statistics ($\bar{Y} = 0.5$, $S^2 = 0.2^2$) are represented by the green triangle; the released statistics ($\bar{Y}^* = 0.43$, $S^2 = 0.28^2$) are represented by the red circle; and, the analyst's posterior mode is represented by the blue diamond. Shaded area represents A , the feasible region for (μ, σ^2) . Plot based on 25,000 Gibbs iterations, thinned to every fifth draw.

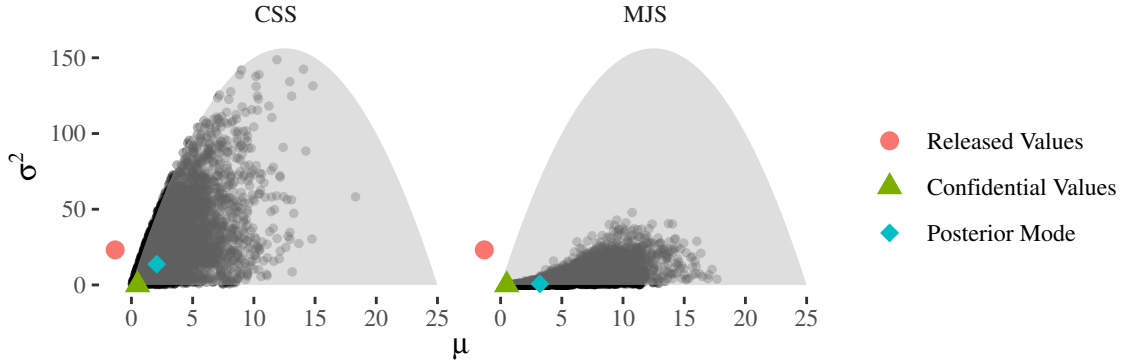


Figure 12: Repeating analysis in Figure 11 but with $(\bar{Y}^* = -1.3, S^2 = 4.8^2)$.

with high density has the strong effect of favoring (μ, σ^2) values that make it less likely to generate Y_i near the boundaries. Additionally, as discussed in Section 4.1, the constrained uniform prior on (μ, σ^2) induces a marginal prior for μ with more probability mass in the center of the feasible region, heightening the discrepancy between the observed $(\bar{Y}^*, S^{2*})$ and the posterior inference for (μ, σ^2) .

I.1.4. Coverage and Error Comparison. The previous section considered only two examples to provide intuition. We now examine the performance of MJS under repeated sampling. For context, we also present results for CSS. We use the simulation study described in Section 4.1 with $\varepsilon = 0.2$; the results are presented in Figure 13. Once n is large enough (e.g., $n \geq 100$), we see little difference in the average interval lengths or RMSEs of the point estimators

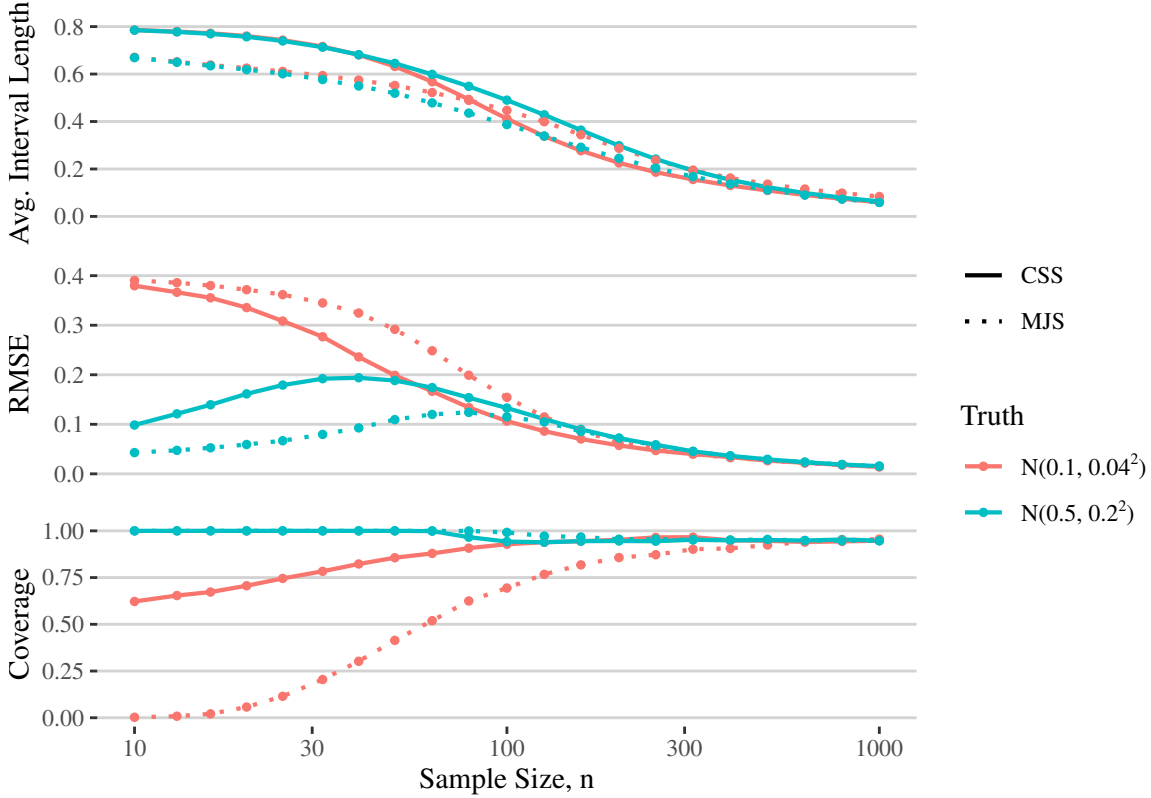


Figure 13: Average 95% HPD interval length for μ (top), root mean square error for estimating μ (middle), and empirical coverage rate of 95% HPD intervals for μ (bottom) for different sample sizes. Analyses with CSS are solid lines and with MJS are dotted lines. The data generating model is either $\mathcal{N}(\mu = 0.1, \sigma^2 = 0.04^2)$ (red) or $\mathcal{N}(\mu = 0.5, \sigma^2 = 0.2^2)$ (blue). DP implemented by Laplace mechanisms with $\varepsilon_1 = \varepsilon_2 = 0.1$ and bounding $Y_i \in [0, 1]$. Each Gibbs sampler is run for 20,000 iterations. Results based on 10,000 simulated datasets.

for the two methods. Until n reaches 500, the confidence interval coverage rate for MJS is lower than 95% when the parameters are near the boundary. The coverage rates for CSS are nearly 95% for all values of $n \geq 100$, with similar average lengths and RMSEs as MJS.

When $n < 100$, we see differences in the performances of the two methods, especially with respect to the coverage rates. When the true μ is near the center, both methods have near 100% coverage for small n . When the true μ is near the boundary, MJS produces interval estimates with low coverage for small n . The coverage is below 70% for $n < 100$ and near zero when n is very small. CSS offers higher coverage rates, although they remain below nominal when $n < 100$. We also see that, for the smaller values of n , MJS has shorter average interval lengths than CSS, reflecting its tendency to favor small values of σ^2 as discussed in Appendix I.1.3. Additionally, for μ near the boundary of the parameter space and small n , MJS tends to have larger RMSE, reflecting the tendency of the sampler (and

the prior) to favor central values of μ as discussed in Appendix I.1.3. This natural tendency is a benefit when $\mu = 0.5$, but a detriment when $\mu = 0.1$.

In practice, of course, an analyst does not know whether μ is near the center of the region or near the boundary. One possibility is to spend some privacy budget to assess which model may be preferred. For example, the data curator may provide some differentially private measure of how many observations are close to the bounds, which analysts might be able to use to decide which of the models is likely to be more accurate. Additionally, analysts might consider utilizing a different default prior distribution, such as one where the marginal distribution for μ is uniform. Development of these ideas is an interesting avenue for future work.

The compute time required to run the MJS sampler on 10,000 simulated datasets required less than 24 hours for $n \leq 50$, 5 days for $n = 500$, and 10 days for $n = 1,000$. The increased compute time compared to what was reported in Appendix J is due primarily to repeatedly running the MJS sampler, which has computational complexity $\mathcal{O}(n)$ per Gibbs iteration.

APPENDIX J. COMPUTE DETAILS

In this section, we describe information on the computer resources needed to reproduce the experiments. We note that Figures 1, 4, 5, and 7 can be run on a personal computer in a few seconds and do not require any additional computing resources. Figure 6 can be run on a personal computer in under five minutes and so also does not require any additional computing resources.

Figures 2 and 3 are based on a large simulation study run on an internal shared compute cluster. The cluster was primarily used for parallelization of computations over a large number of cores; no substantial memory or storage was required. Simulations were run for two values of ϵ , two data generating distributions, and 21 sample sizes for each of the unconstrained and constrained analysis. This yielded 168 input combinations, each of which was run on a different core. The compute time required to run the Gibbs sampler on 10,000 simulated datasets varied for each input combination, but was generally between 10 and 20 hours.