

Locally Private Bayesian Inference for Count Models

Aaron Schein¹ Zhiwei Steven Wu² Alexandra Schofield³ Mingyuan Zhou⁴ Hanna Wallach²

Abstract

We present a general method for privacy-preserving Bayesian inference in Poisson factorization, a broad class of models that includes some of the most widely used models in the social sciences. Our method satisfies limited precision local privacy, a generalization of local differential privacy, which we introduce to formulate privacy guarantees appropriate for sparse count data. We develop an MCMC algorithm that approximates the locally private posterior over model parameters given data that has been locally privatized by the geometric mechanism (Ghosh et al., 2012). Our solution is based on two insights: 1) a novel reinterpretation of the geometric mechanism in terms of the Skellam distribution (Skellam, 1946) and 2) a general theorem that relates the Skellam to the Bessel distribution (Yuan & Kalbfleisch, 2000). We demonstrate our method in two case studies on real-world email data in which we show that our method consistently outperforms the commonly-used naïve approach, obtaining higher quality topics in text and more accurate link prediction in networks. On some tasks, our privacy-preserving method even outperforms non-private inference which conditions on the true data.

1. Introduction

Data from social processes often take the form of discrete observations (e.g., edges in a social network, word tokens in an email) and these observations often contain sensitive information. As more aspects of social interaction are digitally recorded, the opportunities for social scientific insights grow; however, so too does the risk of unacceptable privacy violations. As a result, there is a growing need to develop privacy-preserving data analysis methods. In practice, social scientists will be more likely to adopt these methods if doing

so entails minimal change to their current methodology.

Toward that end, we present a method for privacy-preserving Bayesian inference in Poisson factorization (Titsias, 2008; Cemgil, 2009; Zhou & Carin, 2012; Gopalan & Blei, 2013; Paisley et al., 2014), a broad class of models for learning latent structure from discrete data. This class contains many of the most widely used models in the social sciences, including topic models for text corpora (Blei et al., 2003; Buntine & Jakulin, 2004; Canny, 2004), population models for genetic data (Pritchard et al., 2000), stochastic block models for social networks (Ball et al., 2011; Gopalan & Blei, 2013; Zhou, 2015), and tensor factorization for dyadic data (Welling & Weber, 2001; Chi & Kolda, 2012; Schmidt & Morup, 2013; Schein et al., 2015; 2016b). It further includes deep hierarchical models (Ranganath et al., 2015; Zhou et al., 2015), dynamic models (Charlin et al., 2015; Acharya et al., 2015; Schein et al., 2016a), and many others.

Our method applies principles of differential privacy (Dwork et al., 2006) to observations locally before they are aggregated into a data set for analysis; this ensures that no single centralized server need ever store the non-privatized data, a condition that is often desirable. We introduce *limited precision local privacy* (LPLP)—which generalizes the standard definition of local differential privacy—in order to formulate privacy guarantees that are natural for applications like topic modeling, in which observations (i.e., documents) are high-dimensional vectors of sparse counts. To privatize count data, our method applies the geometric mechanism (Ghosh et al., 2012) which we prove is a mechanism for LPLP.

Under local privacy, a data analysis algorithm sees only privatized (i.e., noised) versions of data. A key statistical challenge is how to fit a model of the true data having only observed the privatized data. One option is a naïve approach, wherein inference proceeds as usual, treating the privatized data as if it were not privatized. In the context of maximum likelihood estimation, the naïve approach has been shown to exhibit pathologies when observations are discrete or count-valued; researchers have therefore advocated for treating the non-privatized observations as latent variables to be inferred (Yang et al., 2012; Karwa et al., 2014; Bernstein et al., 2017). We embrace this approach and extend it to Bayesian inference, where our goal is to approximate the *locally private posterior* distribution over latent variables conditioned on the privatized data and randomized response method.

¹University of Massachusetts Amherst ²Microsoft Research, New York ³Cornell University ⁴University of Texas at Austin. Correspondence to: Aaron Schein <aschein@cs.umass.edu>.

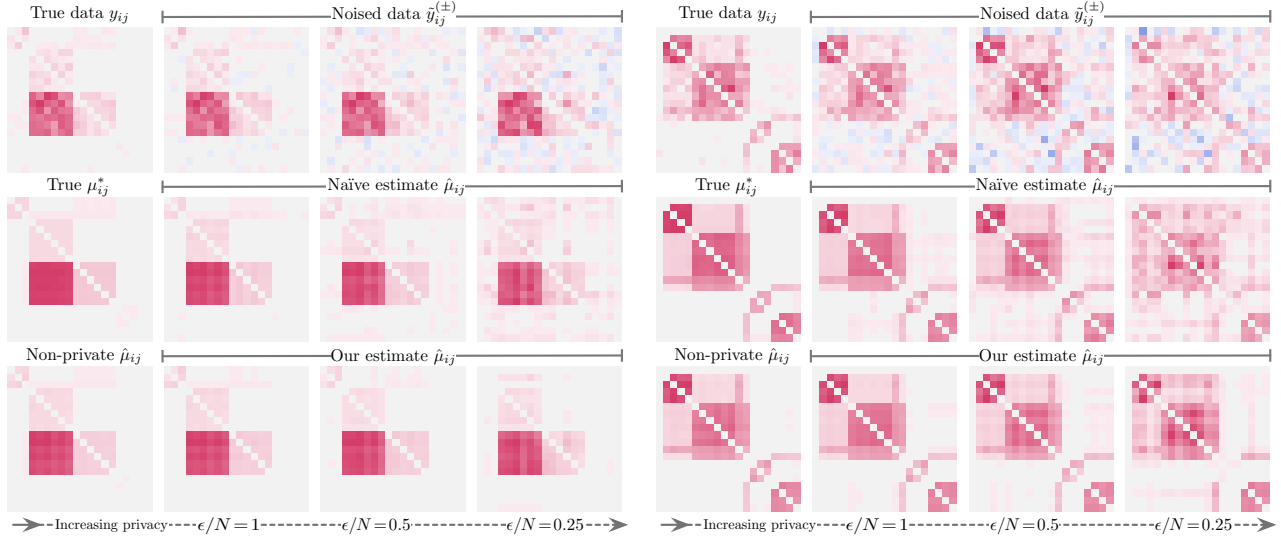


Figure 1. Block structure recovery: our method vs. the naïve approach. We generated the non-privatized data synthetically. We then privatized the data using three different levels of noise. The top row depicts the data, using red to denote positive observations and blue to denote negative observations. As privacy increases, the naïve approach overfits the noise and fails to recover the true μ_{ij}^* values, predicting high values even for sparse parts of the matrix. In contrast, our method recovers the latent structure, even at high noise levels.

Towards that goal, we introduce a novel Markov chain Monte Carlo (MCMC) algorithm that is asymptotically guaranteed to sample from the locally private posterior. Our algorithm iteratively re-samples values of the true underlying data from its inverse distribution conditioned on the model latent variables, and then vice versa; as a result our algorithm is modular, allowing a user to re-purpose an existing implementation of non-private MCMC which samples the model latent variables given (samples of) the true data.

Our main technical contribution is the derivation of an analytically closed-form and computationally tractable scheme for “inverting the noise process” to sample values of the true data. Our solution relies two key insights: 1) a novel reinterpretation of the geometric mechanism in terms of the Skellam distribution (Skellam, 1946), and 2) a theorem that relates the Skellam to the Bessel (Yuan & Kalbfleisch, 2000) distribution. These insights are generally applicable and may be of independent interest in other contexts.

We present two case studies applying our method to 1) topic modeling for text corpora and, 2) overlapping community detection in social networks. We compare our method to the naïve approach over a range of privacy levels on synthetic data and real-world data from the Enron email corpus (Klimt & Yang, 2004). Our method consistently outperforms the naïve approach on a range of intrinsic and extrinsic evaluation metrics—notably, our method obtains more coherent topics from text and more accurate predictions of missing links in networks. We provide an illustrative example in figure 1. Finally, on some tasks, our method even outperforms the non-private method which conditions on the true data.

This finding suggests that Poisson factorization, as widely deployed in practice, is under-regularized and prompts future research into applying and understanding our method as a general method for robust count modeling.

2. Background and problem formulation

Differential privacy. Differential privacy (Dwork et al., 2006) is a rigorous privacy criterion that guarantees that no single observation in a data set will have a significant influence on the information obtained by analyzing that data set.

Definition 1. A randomized algorithm $\mathcal{A}(\cdot)$ satisfies ϵ -differential privacy if for all pairs of neighboring data sets Y and Y' that differ in only a single observation

$$P(\mathcal{A}(Y) \in \mathcal{S}) \leq e^\epsilon P(\mathcal{A}(Y') \in \mathcal{S}), \quad (1)$$

for all subsets $\mathcal{S}$ in the range of $\mathcal{A}(\cdot)$.

Local differential privacy. We focus on local differential privacy, which we refer to as local privacy. Under this criterion, the observations remain private from even the data analysis algorithm. The algorithm only sees privatized versions of the observations, often constructed by adding noise from specific distributions. The process of adding noise is known as randomized response—a reference to survey-sampling methods originally developed in the social sciences prior to the development of differential privacy (Warner, 1965). Satisfying this criterion means one never needs to aggregate the true data in a single location.

Definition 2. A randomized response method $\mathcal{R}(\cdot)$ is ϵ -private if for all pairs of observations $y, y' \in \mathcal{Y}$

$$P(\mathcal{R}(y) \in \mathcal{S}) \leq e^\epsilon P(\mathcal{R}(y') \in \mathcal{S}), \quad (2)$$

for all subsets $\mathcal{S}$ in the range of $\mathcal{R}(\cdot)$. If a data analysis algorithm sees only the observations' ϵ -private responses, then the data analysis itself satisfies ϵ -local privacy.

The meaning of ‘‘observation’’ in definitions 1 and 2 varies across applications. In the context of topic modeling, an observation is an individual document $y \in \mathbb{Z}_+^V$ such that each single entry $y_v \in \mathbb{Z}_+$ is the count of word tokens of type v in the document. To guarantee local privacy, the randomized response method has to satisfy the condition in equation 2 for any pair of observations. This typically involves adding noise that scales with the maximum difference between any pair of observations defined as $N^{(\max)} = \max_{y, y'} \|y - y'\|_1$. When the observations are documents, $N^{(\max)}$ can be prohibitively large and the amount of noise will overwhelm the signal in the data. This motivates the following alternative formulation of privacy.

Limited-precision local privacy. While standard local privacy requires that arbitrarily different observations become indistinguishable under the randomized response method, this guarantee may be unnecessarily strong in some settings. For instance, suppose a user would like to hide only the fact that their email contains a handful of vulgar curse words. Then it is sufficient to have a randomized response method which guarantees that any two similar emails—one containing the vulgar curse words and the same email without them—will be rendered indistinguishable after randomization. In other words, this only requires the randomized response method to render small neighborhoods of possible observations indistinguishable. To operationalize this kind of guarantee, we generalize definition 2 and define limited-precision local privacy (LPLP).

Definition 3. For any positive integer N , we say that a randomized response method $\mathcal{R}(\cdot)$ is (N, ϵ) -private if for all pairs of observations $y, y' \in \mathcal{Y}$ such that $\|y - y'\|_1 \leq N$

$$P(\mathcal{R}(y) \in \mathcal{S}) \leq e^\epsilon P(\mathcal{R}(y') \in \mathcal{S}), \quad (3)$$

for all subsets $\mathcal{S}$ in the range of $\mathcal{R}(\cdot)$. If a data analysis algorithm sees only the observations' (N, ϵ) -private responses, then the data analysis itself satisfies (N, ϵ) -limited-precision local privacy. If $\|y\|_1 \leq N$ for all $y \in \mathcal{Y}$, then (N, ϵ) -limited-precision local privacy implies ϵ -local privacy.

LPLP is the local privacy analog to limited-precision differential privacy, originally proposed by Flood et al. (2013) and subsequently used to privatize analyses of geographic location data (Andrés et al., 2013) and financial network data (Papadimitriou et al., 2017). We emphasize that this is a *strict generalization* of local privacy. A

randomized response method that satisfies LPLP adds noise which scales as a function of N and ϵ —thus the same method may be interpreted as being ϵ -private for a given setting of N or ϵ' -private for a different setting N' . In section 3, we describe the geometric mechanism (Ghosh et al., 2012) and show how it satisfies LPLP.

Differentially private Bayesian inference. In Bayesian statistics, we begin with a probabilistic model $\mathcal{M}$ that relates observable variables Y to latent variables Z via a joint distribution $P_{\mathcal{M}}(Y, Z)$. The goal of inference is then to compute the posterior distribution $P_{\mathcal{M}}(Z | Y)$ over the latent variables conditioned on observed values of Y . The posterior is almost always analytically intractable and thus inference involves approximating it. The two most common methods of approximate Bayesian inference are variational inference, wherein we fit the parameters of an approximating distribution $Q(Z | Y)$, and Markov chain Monte Carlo (MCMC), wherein we approximate the posterior with a finite set of samples $\{Z^{(s)}\}_{s=1}^S$ generated via a Markov chain whose stationary distribution is the exact posterior.

We can conceptualize Bayesian inference as a randomized algorithm $\mathcal{A}(\cdot)$ that returns an approximation to the posterior distribution $P_{\mathcal{M}}(Z | Y)$. In general $\mathcal{A}(\cdot)$ does not satisfy ϵ -differential privacy. However, if $\mathcal{A}(\cdot)$ is an MCMC algorithm that returns a single sample from the posterior, it guarantees privacy (Dimitrakakis et al., 2014; Wang et al., 2015; Foulds et al., 2016; Dimitrakakis et al., 2017). Adding noise to posterior samples can also guarantee privacy (Zhang et al., 2016), though this set of noised samples $\{\tilde{Z}^{(s)}\}_{s=1}^S$ approximates some distribution $\tilde{P}_{\mathcal{M}}(Z | Y)$ that depends on ϵ and is different than the exact posterior (but close, in some sense, and equal when $\epsilon \rightarrow 0$). For specific models, we can also noise the transition kernel of the MCMC algorithm to construct a Markov chain whose stationary distribution is again not the exact posterior, but something close that guarantees privacy (Foulds et al., 2016). We can also take an analogous approach to privatize variational inference, wherein we add noise to the sufficient statistics computed in each iteration (Park et al., 2016).

Locally private Bayesian inference. We first formalize the general objective of Bayesian inference under local privacy. Given a generative model $\mathcal{M}$ for non-privatized data Y and latent variables Z with joint distribution $P_{\mathcal{M}}(Y, Z)$, we further assume a randomized response method $\mathcal{R}(\cdot)$ that generates privatized data sets: $\tilde{Y} \sim P_{\mathcal{R}}(\tilde{Y} | Y)$. The goal is then to approximate the *locally private posterior*:

$$\begin{aligned} P_{\mathcal{M}, \mathcal{R}}(Z | \tilde{Y}) &= \mathbb{E}_{P_{\mathcal{M}, \mathcal{R}}(Y | \tilde{Y})} [P_{\mathcal{M}}(Z | Y)] \\ &= \int P_{\mathcal{M}}(Z | Y) P_{\mathcal{M}, \mathcal{R}}(Y | \tilde{Y}) dY. \end{aligned} \quad (4)$$

This distribution correctly characterizes our uncertainty about the latent variables Z , conditioned on all of our

observations and assumptions—i.e., the privatized data $\tilde{Y}$, the model $\mathcal{M}$, and the randomized response method $\mathcal{R}$. The expansion in equation 4 shows that this posterior implicitly treats the non-privatized data set Y as a latent variable and marginalizes over it using the mixing distribution $P_{\mathcal{M},\mathcal{R}}(Y|\tilde{Y})$ which is itself a posterior that characterizes our uncertainty about Y given all we observe. The central point here is that if we can generate samples from $P_{\mathcal{M},\mathcal{R}}(Y|\tilde{Y})$, we can use them to approximate the expectation in equation 4, assuming that we already have a method for approximating the non-private posterior $P_{\mathcal{M}}(Z|Y)$. In the context of MCMC, iteratively re-sampling values of the non-privatized data from their complete conditional—i.e., $Y^{(s)} \sim P_{\mathcal{M},\mathcal{R}}(Y|Z^{(s-1)}, \tilde{Y})$ —and then re-sampling values of the latent variables—i.e., $Z^{(s)} \sim P_{\mathcal{M}}(Z|Y^{(s)})$ —constitutes a Markov chain whose stationary distribution is $P_{\mathcal{M},\mathcal{R}}(Z, Y|\tilde{Y})$. In scenarios where we already have derivations and implementations for sampling from $P_{\mathcal{M}}(Z|Y)$, we need only be able to sample efficiently from $P_{\mathcal{M},\mathcal{R}}(Y|Z, \tilde{Y})$ in order to obtain a locally private Bayesian inference algorithm; whether we can do this efficiently depends on our choice of $\mathcal{M}$ and $\mathcal{R}$.

We note that the objective of Bayesian inference under local privacy, as defined in equation 4, is similar to that of Williams & McSherry (2010), who identify their key barrier to inference as being unable to analytically form the marginal likelihood that links the privatized data to Z :

$$P_{\mathcal{M},\mathcal{R}}(\tilde{Y}|Z) = \int P_{\mathcal{R}}(\tilde{Y}|Y) P_{\mathcal{M}}(Y|Z) dY. \quad (5)$$

In the next sections, we show that if $\mathcal{M}$ is a Poisson factorization model and $\mathcal{R}$ is the geometric mechanism, then we can form an augmented version of this marginal likelihood analytically and derive an MCMC algorithm that samples efficiently from the posterior in equation 4.

3. Locally private Poisson factorization

In this section, we describe a particular model $\mathcal{M}$ —i.e., Poisson factorization—and randomized response method $\mathcal{R}$ —i.e., the geometric mechanism—that are each natural choices for count data. We prove two theorems about $\mathcal{R}$: 1) that it is a mechanism for LPLP and 2) that it can be re-interpreted in terms of the Skellam distribution (Skellam, 1946). We rely on the second theorem to show that our choices of $\mathcal{M}$ and $\mathcal{R}$ combine to yield a novel generative process for privatized count data which we exploit, in the next section, to derive efficient Bayesian inference.

$\mathcal{M}$: Poisson factorization. We assume that Y is a data set consisting of counts, each of which $y_i \in \mathbb{Z}_+$ is an independent Poisson random variable $y_i \sim \text{Pois}(\mu_i)$ with rate μ_i that is defined to be a deterministic function of shared model parameters Z . The subscript i is a multi-index. In Poisson

matrix factorization $\mathbf{i} \equiv (\mathbf{i}_1, \mathbf{i}_2)$. However this formalism includes tensor factorization $\mathbf{i} \equiv (\mathbf{i}_1, \dots, \mathbf{i}_M)$ and/or multiview models in which the multi-index across counts may differ in the number of indices. This class of models includes many widely used models in social science, as described in section 1. In section 5, we provide case studies using two different Poisson factorization models—i.e., mixed-membership stochastic block model for social networks (Ball et al., 2011; Gopalan & Blei, 2013; Zhou, 2015) and latent Dirichlet allocation (Blei et al., 2003)—which we describe in more detail. We selected these models for purposes of exposition—they are simple and widely used. Although both are instances of Poisson matrix factorization, the method we present in this paper applies to any Poisson factorization model.

$\mathcal{R}$: Geometric mechanism. The two most commonly used randomized response mechanisms in the privacy toolbox—the Gaussian and Laplace mechanisms—privatize observations by adding noise drawn from continuous real-valued distributions; they are thus unnatural choices for count data. Ghosh et al. (2012) introduced the geometric mechanism, which can be viewed as the discrete analog to the Laplace mechanism, and involves adding integer-valued noise $\tau \in \mathbb{Z}$ drawn from the two-sided geometric distribution $\tau \sim 2\text{Geo}(\alpha)$; the PMF of this distribution is:

$$2\text{Geo}(\tau; \alpha) = \frac{1 - \alpha}{1 + \alpha} \alpha^{|\tau|}. \quad (6)$$

Theorem 1. (Proof in appendix) *Let randomized response method $\mathcal{R}(\cdot)$ be the geometric mechanism with parameter α . Then for any positive integer N , and any pair of observations $y, y' \in \mathcal{Y}$ such that $\|y - y'\|_1 \leq N$, $\mathcal{R}(\cdot)$ satisfies*

$$P(\mathcal{R}(y) \in \mathcal{S}) \leq e^\epsilon P(\mathcal{R}(y') \in \mathcal{S}), \quad (7)$$

for all subsets $\mathcal{S}$ in the range of $\mathcal{R}(\cdot)$, where

$$\epsilon = N \ln \left(\frac{1}{\alpha} \right). \quad (8)$$

Therefore, for any N , the geometric mechanism with parameter α is an (N, ϵ) -private randomized response method with $\epsilon = N \ln \left(\frac{1}{\alpha} \right)$. If a data analysis algorithm sees only the observations' (N, ϵ) -private responses, then the data analysis satisfies (N, ϵ) -limited precision local privacy.

Theorem 2. (Proof in appendix) *A two-sided geometric random variable $\tau \sim 2\text{Geo}(\alpha)$ can be generated as follows:*

$$\tau \sim \text{Skel}(\lambda^{(+)}, \lambda^{(-)}), \quad \lambda^{(*)} \sim \text{Exp}\left(\frac{\alpha}{1-\alpha}\right), \quad (9)$$

where $\text{Exp}(\cdot)$ is the exponential distribution and the Skellam distribution is the marginal distribution over the difference $\tau := g^{(+)} - g^{(-)}$ of two independent Poisson random variables $g^{(*)} \sim \text{Pois}(\lambda^{(*)})$ for $* \in \{+, -\}$.

Combining $\mathcal{M}$ and $\mathcal{R}$. We assume each true count is generated by $\mathcal{M}$ —i.e., $y_i \sim \text{Pois}(\mu_i)$ —then privatized by $\mathcal{R}$:

$$\tilde{y}_i^{(\pm)} := y_i + \tau_i, \quad \tau_i \sim 2\text{Geo}(\alpha). \quad (10)$$

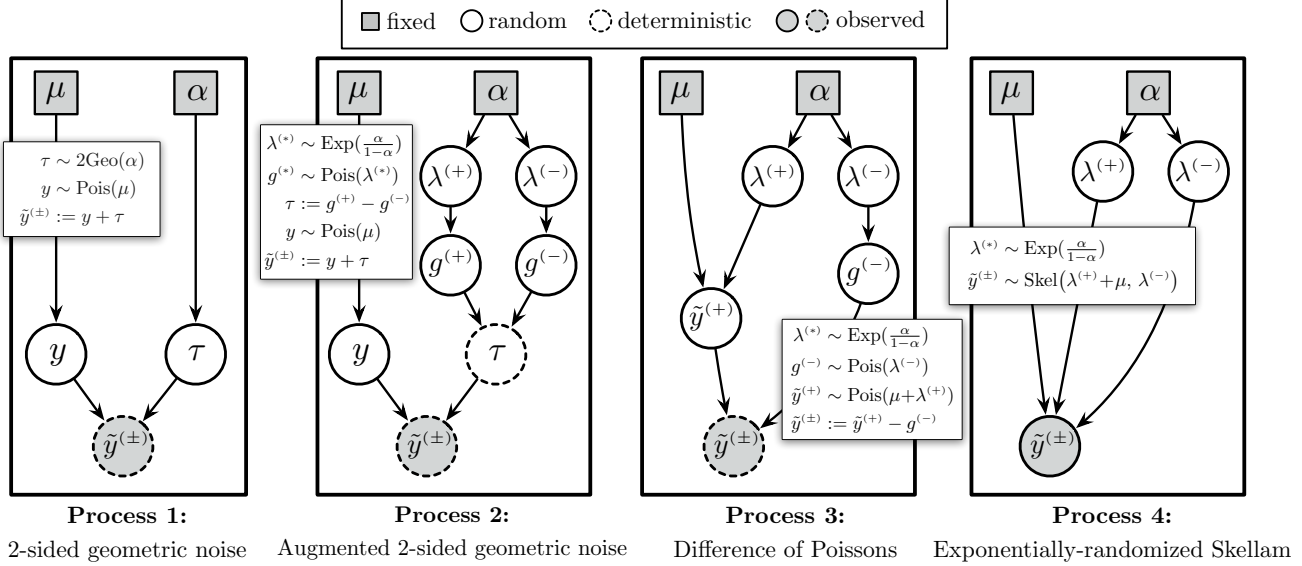


Figure 2. Four generative processes that yield the same marginal distributions $P(\tilde{y}^{(\pm)} | \mu, \alpha)$. Process 1 generates $y^{(\pm)}$ as the sum of an independent Poisson and two-sided geometric random variable. Process 2 augments the two-sided geometric random variable as the difference of two Poisson random variables with exponentially-randomized rates. Process 3 represents the sum of y and the additive geometric random variable $g^{(+)}$ as a single Poisson random variable $\tilde{y}^{(+)}$. Process 4 marginalizes out the remaining Poisson random variables to yield a generative process for $\tilde{y}^{(\pm)}$ as a Skellam random variable with exponentially-randomized rates.

where $\tilde{y}_i^{(\pm)}$ is the privatized observation which we superscript with $(\pm)$ to denote that it may be non-negative or negative since the additive noise $\tau_i \in \mathbb{Z}$ may be negative.

Via theorem 2, we can express the generative process for $\tilde{y}_i^{(\pm)}$ in four equivalent ways, shown in figure 2, each of which provides a unique and necessary insight. Process 1 is a graphical representation of the generative process as defined thus far. Process 2 represents the two-sided geometric noise in terms of a pair of Poisson random variables with exponentially distributed rates; in so doing, it reveals the auxiliary variables that facilitate inference. Process 3 represents the sum of the true count and the positive component of the noise as a single Poisson random variable $\tilde{y}_i^{(+)} = y_i + g_i^{(+)}$. Process 4 marginalizes out the remaining Poisson random variables to obtain a marginal representation $\tilde{y}_i^{(\pm)}$ as a Skellam random variable with exponentially-randomized rates:

$$\tilde{y}_i^{(\pm)} \sim \text{Skel}(\lambda_i^{(+)} + \mu_i, \lambda_i^{(-)}), \quad \lambda_i^{(*)} \sim \text{Exp}\left(\frac{\alpha}{1-\alpha}\right). \quad (11)$$

The derivation of processes 2–4 relies on properties of the two-sided geometric, Skellam, and Poisson distributions which we provide in the appendix.

4. MCMC algorithm

Upon observing a privatized data set $\tilde{Y}^{(\pm)}$, the goal of a Bayesian agent is to approximate the locally private posterior. As explained in section 2, to do so with MCMC, we need only be able to sample the true data y_i as a latent variable from its complete conditional

$P_{\mathcal{M}, \mathcal{R}}(y_i | \tilde{y}_i^{(\pm)}, \mu_i, -)$. By assuming that the privatized observations $\tilde{y}_i^{(\pm)}$ are conditionally independent Skellam random variables, as in equation 11, and we may exploit the following general theorem that relates the Skellam to the Bessel (Yuan & Kalbfleisch, 2000) distribution.

Theorem 3. (Proof in appendix) *Consider two Poisson random variables $y_1 \sim \text{Pois}(\lambda^{(+)})$ and $y_2 \sim \text{Pois}(\lambda^{(-)})$. Their minimum $m := \min\{y_1, y_2\}$ and their difference $\tau := y_1 - y_2$ are deterministic functions of y_1 and y_2 . However, if not conditioned on y_1 and y_2 , the random variables m and τ can be marginally generated as follows:*

$$\tau \sim \text{Skel}(\lambda^{(+)}, \lambda^{(-)}), \quad m \sim \text{Bes}\left(|\tau|, 2\sqrt{2\lambda^{(+)}\lambda^{(-)}}\right). \quad (12)$$

Theorem 3 means that we can generate two independent Poisson random variables by first generating their difference τ and then their minimum m . Because $\tau = y_1 - y_2$, if τ is positive, then y_2 must be the minimum and thus $y_1 = \tau + m$. In practice, this means that if we only get to observe the difference of two Poisson-distributed counts, we can still “recover” the counts by sampling a Bessel auxiliary variable.

Assuming that $\tilde{y}_i^{(\pm)} \sim \text{Skel}(\lambda_i^{(+)} + \mu_i, \lambda_i^{(-)})$ via theorem 2, we first sample an auxiliary Bessel random variable m_i :

$$(m_i | -) \sim \text{Bes}\left(|\tilde{y}_i^{(\pm)}|, 2\sqrt{(\lambda_i^{(+)} + \mu_i)\lambda_i^{(-)}}\right). \quad (13)$$

Yuan & Kalbfleisch (2000) give details of the Bessel distribution, which can be sampled efficiently (Devroye, 2002).

By theorem 3, m_i represents the minimum of two latent Poisson random variables whose difference equals the observed $\tilde{y}_i^{(\pm)}$; these two latent counts are given explicitly in process 3 of figure 2—i.e., $\tilde{y}_i^{(\pm)} := \tilde{y}_i^{(+)} - g_i^{(-)}$ and thus $m_i = \min\{\tilde{y}_i^{(+)}, g_i^{(-)}\}$. Given a sample of m_i and the observed value of $\tilde{y}_i^{(\pm)}$, we can then compute $\tilde{y}_i^{(+)}, g_i^{(-)}$:

$$\begin{aligned} \tilde{y}_i^{(+)} &:= m_i, & g_i^{(-)} &:= \tilde{y}_i^{(+)} - \tilde{y}_i^{(\pm)} \text{ if } \tilde{y}_i^{(\pm)} \leq 0 \\ g_i^{(-)} &:= m_i, & \tilde{y}_i^{(+)} &:= g_i^{(-)} + \tilde{y}_i^{(\pm)} \text{ otherwise.} \end{aligned} \quad (14)$$

Because $\tilde{y}_i^{(+)} = y_i + g_i^{(-)}$ is itself the sum of two independent Poisson random variables, we can then sample y_i from its conditional posterior, which is a binomial distribution:

$$(y_i | -) \sim \text{Binom}\left(\tilde{y}_i^{(+)}, \frac{\mu_i}{\mu_i + \lambda_i^{(+)}}\right). \quad (15)$$

Equations 13 through 15 sample the true underlying data y_i from $P_{\mathcal{M}, \mathcal{R}}(y_i | \tilde{y}_i^{(\pm)}, \mu_i, \lambda_i)$. We may also re-sample the auxiliary variables $\lambda_i^{(+)}, \lambda_i^{(-)}$ from their complete conditional, which is a gamma distribution, by conjugacy:

$$(\lambda_i^{(*)} | -) \sim \Gamma\left(1 + g_i^{(*)}, \frac{\alpha}{1 - \alpha} + 1\right). \quad (16)$$

Iteratively re-sampling y_i and λ_i constitutes a chain whose stationary distribution over y_i is $P_{\mathcal{M}, \mathcal{R}}(y_i | \tilde{y}_i^{(\pm)}, \mu_i)$, as desired. Conditioned on a sample of the underlying data set Y , we then re-sample the latent variables Z (that define the rates μ_i) from their complete conditionals, which match those in standard non-private Poisson factorization. Equations 13–16 along with non-private complete conditionals for Z thus define a privacy-preserving MCMC algorithm that is asymptotically guaranteed to sample from the locally private posterior $P_{\mathcal{M}, \mathcal{R}}(Z | \tilde{Y}^{(\pm)})$ for any Poisson factorization model.

5. Case studies

We present two case studies applying the proposed method to 1) overlapping community detection in social networks and 2) topic modeling for text corpora. In each, we formulate natural local-privacy guarantees, ground them in examples, and demonstrate our method on real and synthetic data.

Enron corpus data. For our real data experiments, we obtained count matrices derived from the Enron email corpus (Klimt & Yang, 2004). For the community detection case study, we obtained a $V \times V$ adjacency matrix Y where y_{ij} is the number of emails sent from actor i to actor j . We included an actor if they sent at least one email and sent or received at least one hundred emails, yielding $V = 161$ actors. When an email included multiple recipients, we incremented the corresponding counts by one. For the topic modeling case study, we randomly selected $D = 10,000$ emails with at least 50 word tokens. We limit the vocabulary to $V = 10,000$ word types, selecting only the most frequent

word types with document frequency less than 0.3. In both case studies, we privatize the data using the geometric mechanism under varying degrees of privacy and examine each method’s ability to reconstruct the true underlying data.

Reference methods. We compare the performance of our method to two references methods: 1) the non-private approach—i.e., standard Poisson factorization fit to the true underlying data, and 2) the naïve approach—i.e., standard Poisson factorization fit to the privatized data, as if it were the true data. The naïve approach must first truncate any negative counts $\tilde{y}_i^{(\pm)} < 0$ to zero and thus implicitly uses the *truncated* geometric mechanism (Ghosh et al., 2012).

Performance measures. All methods generate a set of S samples of the latent variables using MCMC. We use these samples to approximate the posterior expectation of y_i :

$$\hat{\mu}_i = \frac{1}{S} \sum_{s=1}^S \mu_i^{(s)} \approx \mathbb{E}_{P_{\mathcal{M}, \mathcal{R}}(Z | \tilde{Y}^{(\pm)})} [\mu_i]. \quad (17)$$

We then calculate the mean absolute error (MAE) of $\hat{\mu}_i$ with respect to the true data y_i . In the topic modeling case study, we also compare the quality of each method’s inferred latent representation using two different standard metrics.

5.1. Case study 1: Topic modeling

Topic models (Blei et al., 2003) are widely used in the social sciences (Ramage et al., 2009; Grimmer & Stewart, 2013; Mohr & Bogdanov, 2013; Roberts et al., 2013) for characterizing the high-level thematic structure of text corpora via latent “topics”—i.e., probability distributions over vocabulary items. In many settings, documents contain sensitive information (e.g., emails, survey responses) and individuals may be unwilling to share their data without formal privacy guarantees, such as those provided by differential privacy.

Limited-precision local privacy. In this scenario, a data set Y is a $D \times V$ count matrix, each element of which $y_{dv} \in \mathbb{Z}_+$ represents the count of word type $v \in [V]$ in document $d \in [D]$. It is natural to consider each document $\mathbf{y}_d \equiv (y_{d1}, \dots, y_{dV})$ to be a single “observation” we seek to privatize. Under LPLP, the precision level N determines the neighborhood of documents within ϵ -level local privacy applies. For instance, if $N = 4$, then two emails—one that contained four instances of a vulgar curse word and the same one that did not $\|\mathbf{y}_d - \mathbf{y}'_d\|_1 = 4$ —would be rendered indistinguishable after privatization, assuming small ϵ .

Poisson factorization. Gamma–Poisson matrix factorization is commonly used for topic modeling. In this model, Y is a $D \times V$ count matrix. We assume each element is drawn $y_{dv} \sim \text{Pois}(\mu_{dv})$ where $\mu_{dv} = \sum_{k=1}^K \theta_{dk} \phi_{kv}$. The factor θ_{dk} represents how much topic k is used in document d , while the factor ϕ_{kv} represents how much word type v is used in topic k . The set of latent variables is thus

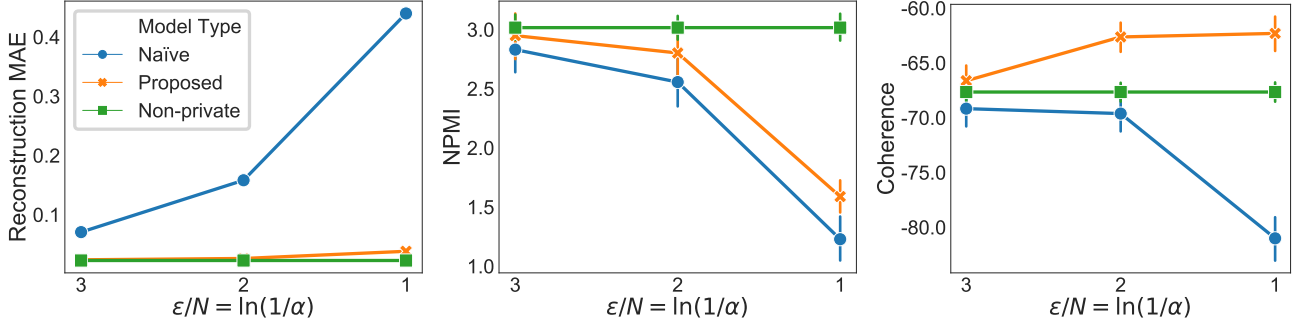


Figure 3. The proposed approach obtains higher quality topics and lower reconstruction error than the naïve approach. Each plot compares the proposed, naïve, and non-private approaches for three increasing levels of noise (privacy) on the Enron corpus data; the non-private values are constant across privacy levels. The *left* plot measures performance using mean absolute error (MAE) of the estimated Poisson rates $\hat{\mu}_{dv}$ against the true underlying data y_{dv} ; *lower is better*. The *center* and *right* plots measure topic quality/interpretability using normalized pointwise mutual information (NPMI) and coherence, respectively, where *higher is better*. When measuring using coherence (*right*), the proposed approach obtains higher quality topics than the non-private approach.

$Z = \{\Theta, \Phi\}$, where Θ and Φ are $D \times K$ and $K \times V$ non-negative, real-valued matrices, respectively. It is standard to assume independent gamma priors over the factors—i.e., $\theta_{dk}, \phi_{kv} \sim \Gamma(a_0, b_0)$, where we set the shape and rate hyperparameters to $a_0 = 0.1$ and $b_0 = 1$, respectively.

Enron corpus experiments. In these experiments, we use the document-by-term count matrix Y derived from the Enron email corpus. We consider three privacy levels $\epsilon/N \in \{3, 2, 1\}$ specified by the ratio of the privacy budget ϵ to the precision N . For each privacy level, we obtain five different privatized matrices, each by adding two-sided geometric noise with $\alpha = -\exp(\epsilon/N)$ independently to each element. We fit both privacy-preserving models—i.e., ours and the naïve approach—to all five privatized matrices for each privacy level. We also fit the non-private approach five independent times to the true matrix. For all models we used $K = 50$ topics. For every model and matrix, we perform 7,500 MCMC iterations, saving every 100th sample of the latent variables $\Phi^{(s)}, \Theta^{(s)}$ after the first 2,500. We use the fifty saved samples to compute $\hat{\mu}_{dv} = \frac{1}{S} \sum_{s=1}^S \sum_{k=1}^K \theta_{dk}^{(s)} \phi_{kv}^{(s)}$.

Results. We find that our approach obtains both lower reconstruction error and higher quality topics than the naïve approach. For each model and matrix, we compute the mean absolute error (MAE) of the reconstruction rates $\hat{\mu}_{dv}$ with respect to the true underlying matrix Y . These results are visualized in the left subplot of figure 3 where we see that the proposed approach reconstructs the true data with nearly as low error as non-private inference (that fits to the true data) while the naïve approach has high error which increases dramatically as the noise (i.e., privacy) increases.

To evaluate topic quality, we use two standard metrics—i.e., normalized pointwise mutual information (NPMI) (Lau et al., 2014) and topic coherence (Mimno et al., 2011)—applied to the 10 words with maximum weight for each sampled topic vector $\phi_k^{(s)}$, using the true data as the refer-

ence corpus. For each method and privacy level, we average these values across samples. The center and right subplots of figure 3 visualize the NPMI and coherence results, respectively. The proposed approach obtains higher quality topics than the naïve approach, as measured by both metrics. As measured by coherence, the proposed approach even obtains higher quality topics than the non-private approach.

Synthetic data experiments. We also find that the proposed approach is more effective than the naïve approach at recovering the ground-truth topics Φ^* from synthetically-generated data. For space reasons, we include these results along with a description of the experiments in the appendix.

5.2. Case study 2: Overlapping community detection

Organizations often ask: are there missing connections between employees that, if present, would significantly reduce duplication of effort? Though social scientists may be able to draw insights based on employees’ interactions, sharing such data risks privacy violations. Moreover, standard anonymization procedures can be reverse-engineered adversarially and thus do not provide privacy guarantees (Narayanan & Shmatikov, 2009). In contrast, the formal privacy guarantees provided by differential privacy may be sufficient for employees to consent to sharing their data.

Limited-precision local privacy. In this setting, a data set Y is a $V \times V$ count matrix, where each element $y_{ij} \in \mathbb{Z}_+$ represents the number of interactions from actor $i \in [V]$ to actor $j \in [V]$. It is natural to consider each count y_{ij} to be a single “observation”. Using the geometric mechanism with $\alpha = -\exp(\epsilon/N)$, if i interacted with j three times $y_{ij} = 3$ and $N = 3$, then an adversary would be unable to tell from $\tilde{y}_{ij}^{(\pm)}$ whether i had interacted with j at all, provided ϵ is sufficiently small. Note that if $y_{ij} \gg N$, then an adversary would be able to tell that i had interacted with j , though not the exact number of times.

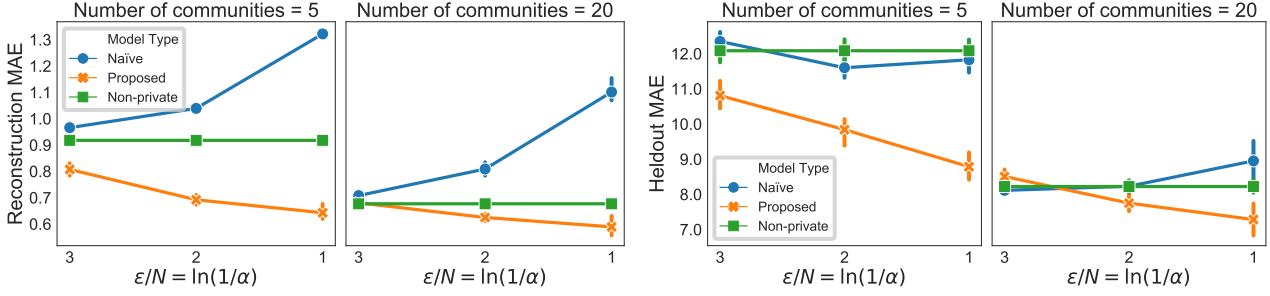


Figure 4. The proposed approach obtains lower error on both reconstruction and heldout link prediction than the naïve and even the non-private approach. *Left and center left:* The methods estimate $\hat{\mu}_{ij}$ from the fully observed matrix; MAE is computed with respect to only the heldout entries y_{ij} . *Center right and right:* The methods estimate $\hat{\mu}_{ij}$ from a partially observed matrix; MAE is computed with respect to only the heldout entries y_{ij} . All experiments are repeated for different numbers of communities $C \in \{5, 10, 20\}$; the results are similar across all values (we include $C = 10$ in the appendix).

Poisson factorization model. The mixed-membership stochastic block model for learning latent overlapping community structure in social networks (Ball et al., 2011; Gopalan & Blei, 2013; Zhou, 2015) is a special case of Poisson factorization where Y is a $V \times V$ count matrix, each element of which is assumed to be drawn $y_{ij} \sim \text{Pois}(\mu_{ij})$ where $\mu_{ij} = \sum_{c=1}^C \sum_{d=1}^C \theta_{ic} \theta_{jd} \pi_{cd}$. The factors θ_{ic} and θ_{jd} represent how much actors i and j participate in communities c and d , respectively, while the factor π_{cd} represents how much actors in community c interact with actors in community d . The set of latent variables is thus $Z = \{\Theta, \Pi\}$ where Θ and Π are $V \times C$ and $C \times C$ non-negative, real-valued matrices, respectively. It is standard to assume independent gamma priors over the factors—i.e., $\theta_{ic}, \pi_{cd} \sim \Gamma(a_0, b_0)$, where we set the shape and rate hyperparameters to $a_0 = 0.1$ and $b_0 = 1$, respectively.

Synthetic data experiments. We generated social networks of $V = 20$ actors with $C = 5$ communities. We randomly generated the true parameters $\theta_{ic}^*, \pi_{cd}^* \sim \Gamma(a_0, b_0)$ with $a_0 = 0.01$ and $b_0 = 0.5$ to encourage sparsity; doing so exaggerates the block structure in the network. We then sampled a data set $y_{ij} \sim \text{Pois}(\mu_{ij}^*)$ and added noise $\tau_{ij} \sim 2\text{Geo}(\alpha)$ for three increasing values of α . In each trial, we set N to the empirical mean of the data $N := \mathbb{E}[y_{ij}]$ and then set $\alpha := \exp(-\epsilon/N)$ for $\epsilon \in \{2.5, 1, 0.75\}$. For each model, we ran 8,500 MCMC iterations, saving every 25th sample after the first 1,000 and using these samples to compute $\hat{\mu}_{ij}$. In figure 1, we visually compare the estimates of $\hat{\mu}_{ij}$ by our proposed method to those of the naïve and non-private approaches. The naïve approach overfits the noise, predicting high rates in sparse parts of the matrix. In contrast, the proposed approach reproduces the sparse block structure even under high noise.

Enron corpus experiments. For the Enron data experiments, we follow the same experimental design outlined in the topic modeling case study; we repeat this experiment using three different numbers of communities $C \in \{5, 10, 20\}$. Each method is applied to five privatized ma-

trices for three different privacy levels. We compute $\hat{\mu}_{ij} = \frac{1}{S} \sum_{s=1}^S \sum_{c=1}^C \sum_{d=1}^C \theta_{ic}^{(s)} \theta_{jd}^{(s)} \pi_{cd}^{(s)}$ from each run and measure reconstruction MAE with respect to the true underlying data y_{ij} . In these experiments, each method observes the entire matrix ($\tilde{Y}^{(\pm)}$ for the privacy-preserving methods and Y for the non-private method). Since missing link prediction is a common task in the networks community, we additionally run the same experiments but where a portion of the matrix is masked—specifically, we hold out all entries $\tilde{y}_{ij}^{(\pm)}$ (or y_{ij} for non-private) that involve any of the top 50 senders i or recipients j . We then compute $\hat{\mu}_{ij}$, as before, but only for the missing entries and compare heldout MAE across methods.

Results. We visualize the results for $C \in \{5, 20\}$ in figure 4. The results for $C = 10$ are similar and given in the appendix. When reconstructing y_{ij} from observed $\tilde{y}_{ij}^{(\pm)}$ (or y_{ij} for non-private), the proposed approach achieves error lower than the naïve approach and lower error than the non-private approach (which directly observes y_{ij}). Similarly, when predicting missing y_{ij} , the proposed approach achieves the lowest error in most settings.

6. Discussion and future work

The proposed privacy-preserving MCMC method for Poisson factorization improves substantially over the commonly-used naïve approach. A surprising finding is that the proposed method was also often better at predicting the true y_{ij} from privatized $\tilde{y}_{ij}^{(\pm)}$ than even the non-private approach. Similarly, the proposed approach inferred more coherent topics. These empirical findings are in fact consistent with known connections between privacy-preserving mechanisms and regularization (Chaudhuri & Monteleoni, 2009). The proposed approach is able to explain natural dispersion in the true data as coming from the randomized response mechanism; it may thus be more robust—i.e., less susceptible to inferring spurious structure—than non-private Poisson factorization. Future application of the model $y_i \sim \text{Skel}(\lambda_i^{(+)} + \mu_i, \lambda_i^{(-)})$ as a robust alternative to Poisson factorization is thus motivated, as is a theoretical characterization of its regularizing properties.

References

- Acharya, A., Ghosh, J., and Zhou, M. Nonparametric Bayesian factor analysis for dynamic count matrices. *arXiv:1512.08996*, 2015.
- Andrés, M. E., Bordenabe, N. E., Chatzikokolakis, K., and Palamidessi, C. Geo-indistinguishability: Differential privacy for location-based systems. In *Proceedings of the 2013 ACM SIGSAC Conference on Computer & Communications Security, CCS '13*, pp. 901–914, New York, NY, USA, 2013. ACM. ISBN 978-1-4503-2477-9.
- Ball, B., Karrer, B., and Newman, M. E. J. Efficient and principled method for detecting communities in networks. *Physical Review E*, 84(3):036103, 2011.
- Bernstein, G., McKenna, R., Sun, T., Sheldon, D., Hay, M., and Miklau, G. Differentially private learning of undirected graphical models using collective graphical models. *arXiv preprint arXiv:1706.04646*, 2017.
- Blei, D. M., Ng, A. Y., and Jordan, M. I. Latent Dirichlet allocation. *Journal of Machine Learning Research*, 3: 993–1022, 2003.
- Buntine, W. and Jakulin, A. Applying discrete PCA in data analysis. In *Proceedings of the 20th Conference on Uncertainty in Artificial Intelligence*, pp. 59–66, 2004.
- Canny, J. GaP: a factor model for discrete data. In *Proceedings of the 27th Annual International ACM SIGIR Conference on Research and Development in Information Retrieval*, pp. 122–129, 2004.
- Cemgil, A. T. Bayesian inference for nonnegative matrix factorisation models. *Computational Intelligence and Neuroscience*, 2009.
- Charlin, L., Ranganath, R., McInerney, J., and Blei, D. M. Dynamic Poisson factorization. In *Proceedings of the 9th ACM Conference on Recommender Systems*, pp. 155–162, 2015.
- Chaudhuri, K. and Monteleoni, C. Privacy-preserving logistic regression. In *Advances in Neural Information Processing Systems*, pp. 289–296, 2009.
- Chi, E. C. and Kolda, T. G. On tensors, sparsity, and non-negative factorizations. *SIAM Journal on Matrix Analysis and Applications*, 33(4):1272–1299, 2012.
- Devroye, L. Simulating Bessel random variables. *Statistics & Probability Letters*, 57(3):249–257, 2002.
- Dimitrakakis, C., Nelson, B., Mitrokotsa, A., and Rubinstein, B. I. P. Robust and private Bayesian inference. In *International Conference on Algorithmic Learning Theory*, pp. 291–305, 2014.
- Dimitrakakis, C., Nelson, B., Zhang, Z., Mitrokotsa, A., and Rubinstein, B. Differential privacy for bayesian inference through posterior sampling. *Journal of Machine Learning Research*, 18(11):1–39, 2017.
- Dwork, C., McSherry, F., Nissim, K., and Smith, A. Calibrating noise to sensitivity in private data analysis. In *Proceedings of the 3rd Theory of Cryptography Conference*, volume 3876, pp. 265–284, 2006.
- Flood, M., Katz, J., Ong, S., and Smith, A. Cryptography and the economics of supervisory information: Balancing transparency and confidentiality. 2013.
- Foulds, J., Geumlek, J., Welling, M., and Chaudhuri, K. On the theory and practice of privacy-preserving Bayesian data analysis. 2016.
- Ghosh, A., Roughgarden, T., and Sundararajan, M. Universally utility-maximizing privacy mechanisms. *SIAM Journal on Computing*, 41(6):1673–1693, 2012.
- Gopalan, P. K. and Blei, D. M. Efficient discovery of overlapping communities in massive networks. *Proceedings of the National Academy of Sciences*, 110(36):14534–14539, 2013.
- Grimmer, J. and Stewart, B. M. Text as data: The promise and pitfalls of automatic content analysis methods for political texts. *Political Analysis*, pp. 1–31, 2013.
- Karwa, V., Slavković, A. B., and Krivitsky, P. Differentially private exponential random graphs. In *International Conference on Privacy in Statistical Databases*, pp. 143–155. Springer, 2014.
- Klimt, B. and Yang, Y. The enron corpus: A new dataset for email classification research. In *European Conference on Machine Learning*, pp. 217–226. Springer, 2004.
- Lau, J. H., Newman, D., and Baldwin, T. Machine reading tea leaves: Automatically evaluating topic coherence and topic model quality. In *Proceedings of the 14th Conference of the European Chapter of the Association for Computational Linguistics*, pp. 530–539, 2014.
- Mimno, D., Wallach, H. M., Talley, E., Leenders, M., and McCallum, A. Optimizing semantic coherence in topic models. In *Proceedings of the Conference on Empirical Methods in Natural Language Processing*, pp. 262–272. Association for Computational Linguistics, 2011.
- Mohr, J. and Bogdanov, P. (eds.). *Poetics: Topic Models and the Cultural Sciences*, volume 41. 2013.
- Narayanan, A. and Shmatikov, V. De-anonymizing social networks. In *Proceedings of the 2009 30th IEEE Symposium on Security and Privacy*, pp. 173–187, 2009.

- Paisley, J., Blei, D. M., and Jordan, M. I. Bayesian non-negative matrix factorization with stochastic variational inference. In Airoldi, E. M., Blei, D. M., Erosheva, E. A., and Fienberg, S. E. (eds.), *Handbook of Mixed Membership Models and Their Applications*, pp. 203–222. 2014.
- Papadimitriou, A., Narayan, A., and Haeberlen, A. Dstress: Efficient differentially private computations on distributed data. In *Proceedings of the Twelfth European Conference on Computer Systems, EuroSys '17*, pp. 560–574, New York, NY, USA, 2017. ACM. ISBN 978-1-4503-4938-3.
- Park, M., Foulds, J., Chaudhuri, K., and Welling, M. Private topic modeling. *arXiv:1609.04120*, 2016.
- Pritchard, J. K., Stephens, M., and Donnelly, P. Inference of population structure using multilocus genotype data. *Genetics*, 155(2):945–959, 2000.
- Ramage, D., Rosen, E., Chuang, J., Manning, C. D., and McFarland, D. A. Topic modeling for the social sciences. In *Neural Information Processing Systems Workshop on Applications for Topic Models*, 2009.
- Ranganath, R., Tang, L., Charlin, L., and Blei, D. Deep exponential families. In *Proceedings of the 18th International Conference on Artificial Intelligence and Statistics*, pp. 762–771, 2015.
- Roberts, M. E., Stewart, B. M., Tingley, D., and Airoldi, E. M. The structural topic model and applied social science. In *Neural Information Processing Systems Workshop on Topic Models: Computation, Application, and Evaluation*, 2013.
- Schein, A., Paisley, J., Blei, D. M., and Wallach, H. Bayesian Poisson tensor factorization for inferring multilateral relations from sparse dyadic event counts. In *Proceedings of the 21st ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pp. 1045–1054, 2015.
- Schein, A., Wallach, H., and Zhou, M. Poisson–gamma dynamical systems. In *Advances in Neural Information Processing Systems* 29, pp. 5005–5013, 2016a.
- Schein, A., Zhou, M., Blei, D. M., and Wallach, H. Bayesian Poisson Tucker decomposition for learning the structure of international relations. In *Proceedings of the 33rd International Conference on Machine Learning*, 2016b.
- Schmidt, M. N. and Morup, M. Nonparametric Bayesian modeling of complex networks: an introduction. *IEEE Signal Processing Magazine*, 30(3):110–128, 2013.
- Skellam, J. G. The frequency distribution of the difference between two Poisson variates belonging to different populations. *Journal of the Royal Statistical Society, Series A (General)*, 109:296, 1946.
- Titsias, M. K. The infinite gamma–Poisson feature model. In *Advances in Neural Information Processing Systems* 21, pp. 1513–1520, 2008.
- Wang, Y.-X., Fienberg, S., and Smola, A. Privacy for free: posterior sampling and stochastic gradient Monte Carlo. In *Proceedings of the 32nd International Conference on Machine Learning*, pp. 2493–2502, 2015.
- Warner, S. L. Randomized response: a survey technique for eliminating evasive answer bias. *Journal of the American Statistical Association*, 60(309):63–69, 1965.
- Welling, M. and Weber, M. Positive tensor factorization. *Pattern Recognition Letters*, 22(12):1255–1261, 2001.
- Williams, O. and McSherry, F. Probabilistic inference and differential privacy. In *Advances in Neural Information Processing Systems* 23, pp. 2451–2459, 2010.
- Yang, X., Fienberg, S. E., and Rinaldo, A. Differential privacy for protecting multi-dimensional contingency table data: Extensions and applications. *Journal of Privacy and Confidentiality*, 4(1):5, 2012.
- Yuan, L. and Kalbfleisch, J. D. On the Bessel distribution and related problems. *Annals of the Institute of Statistical Mathematics*, 52(3):438–447, 2000.
- Zhang, Z., Rubinstein, B. I. P., and Dimitrakakis, C. On the differential privacy of Bayesian inference. In *Proceedings of the 30th AAAI Conference on Artificial Intelligence*, pp. 2365–2371, 2016.
- Zhou, M. Infinite edge partition models for overlapping community detection and link prediction. In *Proceedings of the 18th Conference on Artificial Intelligence and Statistics*, pp. 1135–1143, 2015.
- Zhou, M. and Carin, L. Augment-and-conquer negative binomial processes. In *Advances in Neural Information Processing Systems* 25, pp. 2546–2554, 2012.
- Zhou, M., Cong, Y., and Chen, B. The Poisson gamma belief network. In *Advances in Neural Information Processing Systems* 28, pp. 3043–3051, 2015.