
A Langevin-like Sampler for Discrete Distributions

Ruqi Zhang¹ Xingchao Liu¹ Qiang Liu¹

Abstract

We propose discrete Langevin proposal (DLP), a simple and scalable gradient-based proposal for sampling complex high-dimensional discrete distributions. In contrast to Gibbs sampling-based methods, DLP is able to update all coordinates in parallel in a single step and the magnitude of changes is controlled by a stepsize. This allows a cheap and efficient exploration in the space of high-dimensional and strongly correlated variables. We prove the efficiency of DLP by showing that the asymptotic bias of the stationary distribution is zero for log-quadratic distributions, and is small for distributions that are close to being log-quadratic. With DLP, we develop several variants of sampling algorithms, including unadjusted, Metropolis-adjusted, stochastic and preconditioned versions. DLP outperforms many popular alternatives on a wide variety of tasks, including Ising models, restricted Boltzmann machines, deep energy-based models, binary neural networks and language generation.

1. Introduction

Discrete variables are ubiquitous in machine learning problems ranging from discrete data such as text (Wang & Cho, 2019; Gu et al., 2018) and genome (Wang et al., 2010), to discrete models such as low-precision neural networks (Courbariaux et al., 2016; Peters & Welling, 2018). As data and models become large-scale and complicated, there is an urgent need for efficient sampling from complex high-dimensional discrete distributions.

Markov Chain Monte Carlo (MCMC) methods are typically used to perform sampling, of which the efficiency is largely affected by the proposal distribution (Brooks et al., 2011). For general discrete distributions, Gibbs sampling is broadly applied, which resamples a variable from its conditional

distribution with the remaining variables fixed. Recently, gradient information has been incorporated in the proposal of Gibbs sampling, leading to a substantial boost to the convergence speed of the sampler in discrete spaces (Grathwohl et al., 2021). However, Gibbs-like proposals often suffer from high-dimensional and highly correlated distributions due to conducting a small update per step. In contrast, proposals in continuous spaces that leverage gradients can usually make large effective moves. One of the most popular methods is the Langevin algorithm (Grenander & Miller, 1994; Roberts & Tweedie, 1996; Roberts & Stramer, 2002), which drives the sampler towards high probability regions following a Langevin diffusion. Due to its simplicity and efficiency, the Langevin algorithm has been widely used for sampling from complicated high-dimensional continuous distributions in machine learning and deep learning tasks (Welling & Teh, 2011; Li et al., 2016; Grathwohl et al., 2019; Song & Ermon, 2019). Its great success makes us ask: *what is the simplest and most natural analogue of the Langevin algorithm in discrete domains?*

In this paper, we develop such a Langevin-like proposal for discrete distributions, which can update many coordinates of the variable based on one gradient computation. By reforming the proposal from the standard Langevin algorithm, we find that it can be easily adapted to discrete spaces and can be cheaply computed in parallel due to coordinatewise factorization. We call this proposal *discrete Langevin proposal* (DLP). Inheriting from the Langevin algorithm, DLP is able to update all coordinates in a single step in parallel and the magnitude of changes is controlled by a stepsize. Using this proposal, we are able to obtain high-quality samples conveniently on a variety of tasks. We summarize our contributions as the following:

- We propose discrete Langevin proposal (DLP), a gradient-based proposal for sampling discrete distributions. DLP is able to update many coordinates in a single step with only one gradient computation.
- We theoretically prove the efficiency of DLP by showing that without a Metropolis-Hastings correction, the asymptotic bias of DLP is zero for log-quadratic distributions, and is small for distributions that are close to being log-quadratic.

¹The University of Texas at Austin. Correspondence to: Ruqi Zhang <ruqiz@utexas.edu>.

- With DLP, we develop several variants of sampling algorithms, including unadjusted, Metropolis-adjusted, stochastic and preconditioned versions, indicating the general applicability of DLP for different scenarios.
- We provide extensive experimental results, including Ising models, restricted Boltzmann machines, deep energy-based models, binary Bayesian neural networks and text generation, to demonstrate the superiority of DLP in general settings.

2. Related Work

Gibbs Sampling-based Methods Gibbs sampling is perhaps the *de facto* method for sampling from general discrete distributions. In each step, it iteratively updates one variable leaving the others unchanged. Updating a block of variables is possible, but typically with an increasing cost along with the increase of the block size. To speed up the convergence of Gibbs sampling in high dimensions, Grathwohl et al. (2021) uses gradient information to choose which coordinate to update and Titsias & Yau (2017) introduces auxiliary variables to trade off the number of updated variables in a block for less computation. However, inheriting from Gibbs sampling, these methods still require a large overhead to make significant changes (e.g. > 5 coordinates) to the configuration in one step.

Locally-Balanced Proposals Based on the information of a local neighborhood of the current position, locally-balanced proposals have been developed for sampling from discrete distributions (Zanella, 2020). Later they have been extended to continuous-time Markov processes (Power & Goldman, 2019) and have been tuned via mutual information (Sansone, 2021). Similar to Gibbs sampling-based methods, this type of proposals is very expensive to construct when the local neighborhood is large, preventing them from making large moves in discrete spaces. A recent work (Sun et al., 2022) explores a larger neighborhood by making a sequence of small movements. However, it still only updates one coordinate per gradient computation and the update has to be done in sequence, while on the contrary, our method can update many coordinates based on one gradient computation in parallel.

Continuous Relaxation Incorporating gradients in the proposal has been a great success in continuous spaces, such as the Langevin algorithm, Hamiltonian Monte Carlo (HMC) (Duane et al., 1987; Neal et al., 2011) and their variants. To take advantage of this success, continuous relaxation is applied which performs sampling in a continuous space by gradient-based methods and then transforms the collected samples to the original discrete space (Pakman & Paninski, 2013; Nishimura et al., 2020; Han et al., 2020;

Zhou, 2020; Jaini et al., 2021; Zhang et al., 2022). The efficiency of continuous relaxation highly depends on the properties of the extended continuous distributions which may be difficult to sample from. As shown in previous work, this type of methods usually does not scale to high dimensional discrete distributions (Grathwohl et al., 2021).

3. Preliminaries

We consider sampling from a target distribution

$$\pi(\theta) = \frac{1}{Z} \exp(U(\theta)), \quad \forall \theta \in \Theta,$$

where θ is a d -dimensional variable, Θ is a finite¹ variable domain, U is the energy function, and Z is the normalizing constant for π to be a distribution. In this paper, we restrict to a *factorized* domain, that is $\Theta = \prod_{i=1}^d \Theta_i$, and mainly consider Θ to be $\{0, 1\}^d$ or $\{0, 1, \dots, S-1\}^d$. Additionally, we assume that U can be extended to a differentiable function in $\mathbb{R}^d$. Many popular models have such natural extensions such as Ising models, Potts models, restricted Boltzmann machines, and (deep) energy-based models.

Langevin Algorithm In continuous spaces, one of the most powerful sampling methods is the Langevin algorithm, which follows a Langevin diffusion to update variables:

$$\theta' = \theta + \frac{\alpha}{2} \nabla U(\theta) + \sqrt{\alpha} \xi, \quad \xi \sim \mathcal{N}(0, I_{d \times d}),$$

where α is a stepsize. The gradient helps the sampler to explore high probability regions efficiently. Generally, computing the gradient and sampling a Gaussian variable can be done cheaply in parallel on CPUs and GPUs. As a result, the Langevin algorithm is especially compelling for complex high-dimensional distributions, and has been extensively used in machine learning and deep learning.

4. Discrete Langevin Proposal

In this section, we propose discrete Langevin proposal (DLP), a simple counterpart of the Langevin algorithm in discrete domains.

At the current position θ , the proposal distribution $q(\cdot|\theta)$ produces the next position to move to. As introduced in Section 3, $q(\cdot|\theta)$ of the Langevin algorithm in continuous spaces can be viewed as a Gaussian distribution with mean $\theta + \alpha/2 \nabla U(\theta)$ and covariance $\alpha I_{d \times d}$. Obviously we could not use this Gaussian proposal in discrete spaces. However, we notice that by explicitly indicating the spaces where

¹We consider finite discrete distributions in the paper. However, our algorithms can be easily extended to infinite distributions. See Appendix C for a discussion.

the normalizing constant is computed over, this proposal is essentially applicable to *any* kind of spaces. Specifically, we write out the variable domain Θ explicitly in the proposal distribution,

$$q(\theta'|\theta) = \frac{\exp\left(-\frac{1}{2\alpha} \left\| \theta' - \theta - \frac{\alpha}{2} \nabla U(\theta) \right\|_2^2\right)}{Z_{\Theta}(\theta)}, \quad (1)$$

where the normalizing constant is integrated (continuous) or summed (discrete) over Θ (we use sum below)

$$Z_{\Theta}(\theta) = \sum_{\theta' \in \Theta} \exp\left(-\frac{1}{2\alpha} \left\| \theta' - \theta - \frac{\alpha}{2} \nabla U(\theta) \right\|_2^2\right).$$

Here, Θ can be any space without affecting q being a valid proposal. As a special case, when $\Theta = \mathbb{R}^d$, it follows that $Z_{\mathbb{R}^d}(\theta) = (2\pi\alpha)^{d/2}$ and recovers the Gaussian proposal in the standard Langevin algorithm. When Θ is a discrete space, we naturally obtain a gradient-based proposal for discrete variables.

Computing the sum over the full space in $Z_{\Theta}(\theta)$ is generally very expensive, for example, the cost is $\mathcal{O}(S^d)$ for $\Theta = \{0, 1, \dots, S-1\}^d$. This is why previous methods often restrict their proposals to a small neighborhood. A key feature of the proposal in Equation (1) is that it can be factorized *coordinatewisely*. To see this, we write Equation (1) as $q(\theta'|\theta) = \prod_{i=1}^d q_i(\theta'_i|\theta)$, where $q_i(\theta'_i|\theta)$ is a simple categorical distribution of form:

$$\text{Categorical}\left(\text{Softmax}\left(\frac{1}{2} \nabla U(\theta)_i (\theta'_i - \theta_i) - \frac{(\theta'_i - \theta_i)^2}{2\alpha}\right)\right), \quad (2)$$

with $\theta'_i \in \Theta_i$ (note that the equation does not contain the term $(\alpha/2 \nabla U(\theta)_i)^2$ because it is independent of θ' and will not affect the softmax result). Combining with coordinate-wise factorized domain Θ , the above proposal enables us to update each coordinate in parallel after computing the gradient $\nabla U(\theta)$. The cost of gradient computation is also $\mathcal{O}(d)$, therefore, the overall cost of constructing this proposal depends linearly rather than exponentially on d . This allows the sampler to explore the full space with the gradient information without paying a prohibitive cost.

We denote the proposal in Equation (2) as *Discrete Langevin Proposal* (DLP). DLP can be used with or without a Metropolis-Hastings (MH) step (Metropolis et al., 1953; Hastings, 1970), which is usually combined with proposals to make the Markov chain reversible. Specifically, after generating the next position θ' from a distribution $q(\cdot|\theta)$, the MH step accepts it with probability

$$\min\left(1, \exp(U(\theta') - U(\theta)) \frac{q(\theta|\theta')}{q(\theta'\|\theta)}\right). \quad (3)$$

By rejecting some of the proposed positions, the Markov chain is guaranteed to converge asymptotically to the target distribution.

We outline the sampling algorithms using DLP in Algorithm 1. We call DLP without the MH step as *discrete unadjusted Langevin algorithm* (DULA) and DLP with the MH step as *discrete Metropolis-adjusted Langevin algorithm* (DMALA). Similar to MALA and ULA in continuous spaces (Grenander & Miller, 1994; Roberts & Stramer, 2002), DMALA contains two gradient computations and two function evaluations and is guaranteed to converge to the target distribution, while DULA may have asymptotic bias, but only requires one gradient computation, which is especially valuable when performing the MH step is expensive such as in large-scale Bayesian inference (Welling & Teh, 2011; Durmus & Moulines, 2019).

Connection to Locally-Balanced Proposals Zanella (2020) has developed a class of locally-balanced proposals that can be used in both discrete and continuous spaces. One of the locally-balanced proposals is defined as

$$r(\theta'|\theta) \propto \exp\left(\frac{1}{2}U(\theta') - \frac{1}{2}U(\theta) - \frac{\|\theta' - \theta\|^2}{2\alpha}\right),$$

where $\theta' \in \Theta$. DLP can be viewed as a first-order Taylor series approximation to $r(\theta'|\theta)$ using

$$U(x) - U(\theta) \approx \nabla U(\theta)^\top (x - \theta), \quad \forall x \in \Theta.$$

Zanella (2020) discussed the connection between their proposals and Metropolis-adjusted Langevin algorithm (MALA) in continuous spaces but did not explore it in discrete spaces. Grathwohl et al. (2021) uses a similar Taylor series approximation for another locally-balanced proposal,

$$r'(\theta'|\theta) \propto \exp\left(\frac{1}{2}U(\theta') - \frac{1}{2}U(\theta)\right), \quad (4)$$

where θ' belongs to a hamming ball centered at θ with window size 1. Like Gibbs sampling, their proposal only updates one coordinate per step. They also propose an extension of their method to update X coordinates per step, but with X times gradient computations. See Appendix D for a discussion on Taylor series approximation for Equation (4) without window sizes.

Beyond previous works, we carefully investigate the properties of the Langevin-like proposal in Equation (2) in discrete spaces, providing both convergence analysis and extensive empirical demonstration. We find this simple approach explores the discrete structure surprisingly well, leading to a substantial improvement on a range of tasks.

Algorithm 1 Samplers with Discrete Langevin Proposal (DULA and DMALA).

given: Stepsize α .

loop

for $i = 1 : d$ **do** {Can be done in parallel}

construct $q_i(\cdot|\theta)$ as in Equation (2)

sample $\theta'_i \sim q_i(\cdot|\theta)$

end for

 ▷ Optionally, do the MH step

compute $q(\theta'|\theta) = \prod_i q_i(\theta'_i|\theta)$
 and $q(\theta|\theta') = \prod_i q_i(\theta_i|\theta')$

set $\theta \leftarrow \theta'$ with probability in Equation (3)

end loop

output: samples $\{\theta_k\}$

5. Convergence Analysis for DULA

In the previous section, we showed that DLP is a convenient gradient-based proposal for discrete distributions. However, the effectiveness of a proposal also depends on how close its underlying stationary distribution is to the target distribution. Because if it is far, even if using the MH step to correct the bias, the acceptance probability will be very low. In this section, we provide an asymptotic convergence analysis for DULA (i.e. the sampler using DLP without the MH step). Specifically, we first prove in Section 5.1 that when the stepsize $\alpha \rightarrow 0$, the asymptotic bias of DULA is zero for log-quadratic distributions, which is defined as

$$\pi(\theta) \propto \exp(\theta^\top W \theta + b^\top \theta), \theta \in \Theta \quad (5)$$

with some constants $W \in \mathbb{R}^{d,d}$ and $b \in \mathbb{R}^d$. Without loss of generality, we assume W is symmetric (otherwise we can replace W with $(W + W^\top)/2$ for the eigendecomposition).

Later in Section 5.2, we extend the result to general distributions where we show the asymptotic bias of DULA is small for distributions that are close to being log-quadratic.

5.1. Convergence on Log-Quadratic Distributions

We consider a log-quadratic distribution $\pi(\theta)$ as defined in Equation (5). This type of distributions appears in common tasks such as Ising models. The following theorem summarizes DULA's asymptotic accuracy for such π .

Theorem 5.1. *If the target distribution π is log-quadratic as defined in Equation (5).*

Then the Markov chain following transition $q(\cdot|\theta)$ in Equation (2) (i.e. DULA)

is reversible with respect to some distribution π_α and π_α converges weakly to π as $\alpha \rightarrow 0$. In particular, let $\lambda_{\min}$ be

the smallest eigenvalue of W , then for any $\alpha > 0$,

$$\|\pi_\alpha - \pi\|_1 \leq Z \cdot \exp\left(-\frac{1}{2\alpha} - \frac{\lambda_{\min}}{2}\right),$$

where Z is the normalizing constant of π .

Theorem 5.1 shows that the asymptotic bias of DULA decreases at a $\mathcal{O}(\exp(-1/(2\alpha)))$ rate which vanishes to zero as the stepsize $\alpha \rightarrow 0$. This is similar to the case of the Langevin algorithm in continuous spaces, where it converges asymptotically when the stepsize goes to zero.

We empirically verify this theorem in Figure 1a. We run DULA with varying stepsizes on a 2 by 2 Ising model. For each stepsize, we run the chain long enough to make sure it converged. The results clearly show that the distance between the stationary distribution of DULA and the target distribution decreases as the stepsize decreases. Moreover, the decreasing speed roughly aligns with a function containing $\exp(-1/(2\alpha))$, which demonstrates the convergence rate with respect to α in Theorem 5.1.

5.2. Convergence on General Distributions

To generalize the convergence result from log-quadratic distributions to general distributions, we first note that the stationary distribution of DULA always exists and is unique since its transition matrix is irreducible (Levin & Peres, 2017). Then we want to make sure the bound is tight in the log-quadratic case and derive one that depends on how much the distribution differs from being log-quadratic. A natural measure of this difference is the distance of ∇U to a linear function. Specifically, we assume that $\exists W \in \mathbb{R}^{d \times d}, b \in \mathbb{R}, \epsilon \in \mathbb{R}^+$, such that

$$\|\nabla U(\theta) - (2W\theta + b)\|_1 \leq \epsilon, \forall \theta \in \Theta. \quad (6)$$

Then we have the following theorem for the asymptotic bias of DULA on general distributions.

Theorem 5.2. *Let π be the target distribution and $\pi'(\theta) = \exp(\theta^\top W \theta + b\theta) / Z'$ be the log-quadratic distribution satisfying the assumption in Equation (6), then the stationary distribution of DULA satisfies*

$$\|\pi_\alpha - \pi\|_1 \leq 2c_1 (\exp(c_2\epsilon) - 1) + Z' \exp\left(-\frac{1}{2\alpha} - \frac{\lambda_{\min}}{2}\right),$$

where c_1 is a constant depending on π' and α ; c_2 is a constant depending on Θ and $\max_{\theta, \theta' \in \Theta} \|\theta' - \theta\|_\infty$.

The first term in the bound captures the bias induced by the deviation of π from being log-quadratic, which decreases in a $\mathcal{O}(\exp(\epsilon))$ rate. The second term captures the bias by using a non-zero stepsize which directly follows from Theorem 5.1. To ensure a satisfying convergence, Theorem 5.2 suggests that we should choose a continuous extension for U of which the gradient is close to a linear function.

Example. To empirically verify Theorem 5.2, we consider a 1-d distribution $\pi(\theta) \propto \exp(a\theta^2 + b\theta + 2\epsilon \sin(\theta\pi/2))$ where $\theta \in \{-1, 1\}$ and $\epsilon \in \mathbb{R}$. The gradient is $\nabla U(\theta) = 2a\theta + b + \epsilon\pi \cos(\theta\pi/2)$ of which the closeness to a linear function in $\mathbb{R}$ is controlled by ϵ . From Figure 1b, we can see that when ϵ decreases, the asymptotic bias of DULA decrease, aligning with Theorem 5.2.

In fact, the example of $\pi(\theta) \propto \exp(2\epsilon \sin(\theta\pi/2))$ with $\theta \in \{-1, 1\}$ can be considered as an counterexample for *any* existing gradient-based proposals in discrete domains. The gradient on $\{-1, 1\}$ is always zero regardless the value of ϵ whereas the target distribution clearly depends on ϵ . This suggests that we should be careful with the choice of the continuous extension of U since some extensions will not provide useful gradient information to guide the exploration. Our Theorem 5.2 provides a guide about how to choose such an extension.

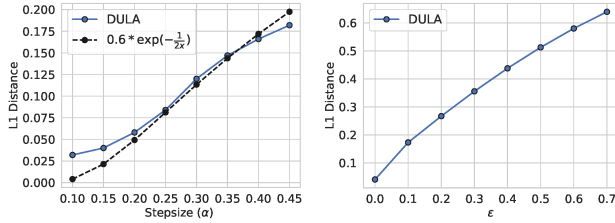


Figure 1. Empirical verification of theorems. **Left:** DULA with varying stepsizes on an Ising model. **Right:** 1-d distribution with varying closeness ϵ to being log-quadratic.

6. Other Variants

Thanks to the similarity with the standard Langevin algorithm, discrete Langevin proposal can be extended to different usage scenarios following the rich literature of the standard Langevin algorithm. We briefly discuss two such variants in this section.

With Stochastic Gradients Similar to Stochastic gradient Langevin dynamics (SGLD) (Welling & Teh, 2011), we can replace the full-batch gradient with an unbiased stochastic estimation $\hat{\nabla}U$ in DLP, which will further reduce the cost of our method on large-scale problems. To show the influence of stochastic estimation, we consider a binary domain for simplicity. We further assume the stochastic gradient has a bounded variance and the norm of the true gradient and the stochastic gradient are bounded. Then we have the following theorem.

Theorem 6.1. Let $\Theta = \{0, 1\}^d$. We assume that the true gradient ∇U and the stochastic gradient $\hat{\nabla}U$ satisfy: $\mathbf{E}[\hat{\nabla}U_i] = \nabla U_i$, $\mathbf{E}[\hat{\nabla}U_i^2] \leq \sigma^2$ for some constant σ ; $|\hat{\nabla}U_i|, |\nabla U_i| \leq L$ for some constant L . Let q_i and $\hat{q}_i$ be the discrete Langevin proposal for the coordinate i using the full-batch gradient and the stochastic gradient respectively,

then

$$\|\mathbf{E}[\hat{q}_i] - q_i\|_1 \leq 2\sigma \cdot \exp\left(-\frac{1}{2\alpha} + L\right).$$

This suggests that when the variance of the stochastic gradient or the stepsize decreases, the stochastic DLP in expectation will be closer to the full-batch DLP. We test DULA with the stochastic gradient on a binary Bayesian neural network and empirically verify it works well in practice in Appendix I.5.

With Preconditioners When π alters more quickly in some coordinates than others, a single stepsize may result in slow mixing. Under this situation, a preconditioner that adapts the stepsize for different coordinates can help alleviate this problem. We show that it is easy for DLP to incorporate diagonal preconditioners, as long as the coordinatewise factorization still holds. For example, when the preconditioner is constant, that is we scale each coordinate by a number g_i ,

discrete Langevin proposal becomes

$$q_i(\theta'_i|\theta) \propto \exp\left(\frac{1}{2}\nabla U(\theta)_i(\theta'_i - \theta_i) - \frac{(\theta_i - \theta'_i)^2}{2\alpha g_i}\right).$$

Similar to continuous spaces, g_i adjusts the stepsize for different coordinates considering their variation speed. The above proposal is obtained by applying a coordinate transformation to θ and then transforming the DLP update back to the original space. In this way, the theoretical results in Section 5 directly apply to it. We put more details and an Ising example in Appendix H.

7. Experiments

We conduct a thorough empirical evaluation of discrete Langevin proposal (DLP), comparing to a range of popular baselines, including Gibbs sampling, Gibbs with Gradient (GWG) (Grathwohl et al., 2021), Hamming ball (HB) (Titsias & Yau, 2017)—three Gibbs-based approaches; discrete Stein Variational Gradient Descent (DSVGD) (Han et al., 2020) and relaxed MALA (RMALA) (Grathwohl et al., 2021)—two continuous relaxation methods; and a locally balanced sampler (LB-1) (Zanella, 2020) which uses the locally-balanced proposal in Equation (4) with window size 1. We denote Gibbs- X for Gibbs sampling with a block-size of X , GWG- X for GWG with X indices being modified per step (see D.2 in Grathwohl et al. (2021) for more details of GWG- X), HB- X - Y for HB with a block size of X and a hamming ball size of Y . All methods are implemented in Pytorch and we use the official release of code from previous papers when possible. In our implementation, DMALA (i.e. discrete Langevin proposal with an MH step)

has a similar cost per step with GWG-1 (the main costs for them are one gradient computation and two function evaluations), which is roughly 2.5x of Gibbs-1. DULA (i.e. discrete Langevin proposal without an MH step) has a similar cost per step with Gibbs-1 (the main cost for DULA is one gradient computation, and for Gibbs-1 is d function evaluations). We released the code at <https://github.com/ruqizhang/discrete-langevin>.

7.1. Sampling From Ising Models

We consider a 5 by 5 lattice Ising model with random variable $\theta \in \{-1, 1\}^d$, and $d = 5 \times 5 = 25$. The energy function is

$$U(\theta) = a\theta^\top W\theta + b\theta,$$

where W is a binary adjacency matrix, $a = 0.1$ is the connectivity strength and $b = 0.2$ is the bias. We first show that DLP can change many coordinates in one iteration while still maintaining a high acceptance rate in Figure 2 (Left). When the stepsize $\alpha = 0.6$, on average DMALA can change 6 coordinates in one iteration with an acceptance rate 52% in the MH step. In comparison, GWG-6 (which at most changes 6 coordinates) only has 43% acceptance rate, not to mention it requires 6x cost of DMALA. This demonstrates that DLP can make large and effective moves in discrete spaces. We compare the root-mean-square error (RMSE) between the estimated mean and the true mean in Figure 3. DMALA is the fastest to converge in terms of both runtime and iterations. This demonstrates the importance of (1) using gradient information to explore the space compared to Gibbs and HB; (2) sampling in the original discrete space compared to DSVGD and RMALA; and (3) changing many coordinates in one step compared to LB-1 and GWG-1. GWG-4 underperforms because of a lower acceptance rate than DMALA. DULA can achieve a similar result as LB-1 and GWG-1 but worse than DMALA, indicating that the MH step accelerates the convergence on this task. In Figure 2 (Right), we compare the effective sample size (ESS) per second for exact samplers (i.e. having the target distribution as its stationary distribution). DMALA significantly outperforms other methods, indicating the correlation among its samples is low due to making significant updates in each step. We additionally present results on Ising models with different connectivity strength a in Appendix I.2.

In what follows, we mainly compare our method with GWG-1 and Gibbs-1, as other methods either could not give reasonable results or are too costly to run.

7.2. Sampling From Restricted Boltzmann Machines

Restricted Boltzmann Machines (RBMs) are generative artificial neural networks, which learn an unnormalized proba-

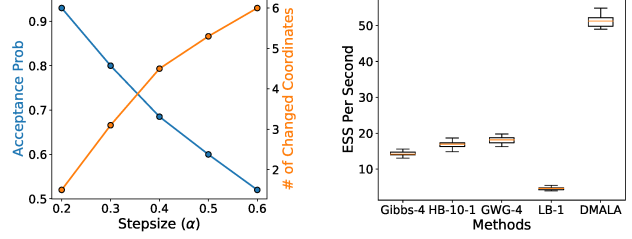


Figure 2. Ising model sampling results. **Left:** DLP is able to change many coordinates while keeping a high acceptance rate. **Right:** DMALA yields the largest effective sample size (ESS) per second among all the methods compared.

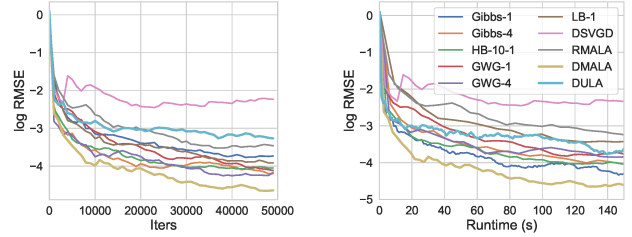


Figure 3. Ising model sampling results. DMALA outperforms the baselines in both number of iterations and running time.

bility over inputs,

$$U(\theta) = \sum_i \text{Softplus}(W\theta + a)_i + b^\top \theta,$$

where $\{W, a, b\}$ are parameters and $\theta \in \{0, 1\}^d$. Following Grathwohl et al. (2021), we train $\{W, a, b\}$ with contrastive divergence (Hinton, 2002) on the MNIST dataset for one epoch. We measure the Maximum Mean Discrepancy (MMD) between the obtained samples and those from Block-Gibbs sampling, which utilizes the known structure.

Results The results are shown in Figure 4. Since DULA and DMALA can update multiple coordinates in a single iteration, they converge remarkably faster than baselines in both iterations and runtime. Furthermore, DMALA reaches the lowest MMD (≈ -6.5) among all the methods after 5,000 iterations, demonstrating the importance of the MH step. We leave the sampled images in Appendix I.3.

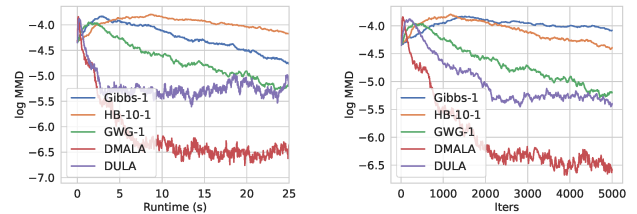


Figure 4. RBM sampling results. DULA and DMALA converge to the true distribution faster than other methods.

Dataset	VAE (Conv)	EBM (Gibbs)	EBM (GWG)	EBM (DULA)	EBM (DMALA)
Static MNIST	-82.41	-117.17	-80.01	-80.71	-79.46
Dynamic MNIST	-80.40	-121.19	-80.51	-81.29	-79.54
Omniglot	-97.65	-142.06	-94.72	-145.68	-91.11
Caltech Silhouettes	-106.35	-163.50	-96.20	-100.52	-87.82

Table 1. EBM learning results. We report the log-likelihoods on the test set for different models trained on discrete image datasets.

Dataset	Training Log-likelihood ($\uparrow$)				Test RMSE ($\downarrow$)			
	Gibbs	GWG	DULA	DMALA	Gibbs	GWG	DULA	DMALA
COMPAS	-0.4063 $\pm$ 0.0016	-0.3729 $\pm$ 0.0019	-0.3590 $\pm$ 0.0003	-0.3394 $\pm$ 0.0016	0.4718 $\pm$ 0.0022	0.4757 $\pm$ 0.0008	0.4879 $\pm$ 0.0024	0.4879 $\pm$ 0.0003
News	-0.2374 $\pm$ 0.0003	-0.2368 $\pm$ 0.0003	-0.2339 $\pm$ 0.0003	-0.2344 $\pm$ 0.0005	0.1048 $\pm$ 0.0025	0.1028 $\pm$ 0.0023	0.0971 $\pm$ 0.0012	0.0943 $\pm$ 0.0009
Adult	-0.4587 $\pm$ 0.0029	-0.4444 $\pm$ 0.0041	-0.3287 $\pm$ 0.0027	-0.3190 $\pm$ 0.0019	0.4826 $\pm$ 0.0027	0.4709 $\pm$ 0.0034	0.3971 $\pm$ 0.0014	0.3931 $\pm$ 0.0019
Blog	-0.3829 $\pm$ 0.0036	-0.3414 $\pm$ 0.0028	-0.2694 $\pm$ 0.0025	-0.2670 $\pm$ 0.0031	0.4191 $\pm$ 0.0061	0.3728 $\pm$ 0.0034	0.3193 $\pm$ 0.0021	0.3198 $\pm$ 0.0043

Table 2. Experiment results with binary Bayesian neural networks on different datasets.

7.3. Learning Energy-based Models

Energy-based models (EBMs) have achieved great success in various areas in machine learning (LeCun et al., 2006). Generally, the density of an energy-based model is defined as $p_\theta(x) = \exp(-E_\theta(x))/Z_\theta$, where E_θ is a function parameterized by θ and Z_θ is the normalizing constant. Training EBM usually involves maximizing the log-likelihood, $\mathcal{L}(\theta) \triangleq \mathbb{E}_{x \sim p_{data}} [\log p_\theta(x)]$. However, direct optimization needs to compute Z_θ , which is intractable in most scenarios. To deal with it, we usually estimate the gradient of the log-likelihood instead,

$$\nabla \mathcal{L}(\theta) = \mathbb{E}_{x \sim p_\theta} [\nabla_\theta E_\theta(x)] - \mathbb{E}_{x \sim p_{data}} [\nabla_\theta E_\theta(x)].$$

Though the first term is easy to estimate from the data, the second term requires samples from p_θ . Better samplers can improve the training process of E_θ , leading to EBMs with higher performance.

7.3.1. ISING MODELS

As in Grathwohl et al. (2021), we generate a 25 by 25 Ising model and generate training data by running a Gibbs sampler. In this experiment, E_θ is an Ising model with learnable parameter $\hat{W}$. We evaluate the samplers by computing RMSE between the estimated $\hat{W}$ and the true W .

Results Our results are summarized in Figure 5. In (a), DMALA and DULA always have smaller RMSE than baselines given the same number of iterations. In (b), DMALA and DULA get a log-RMSE of -5.0 in 800s, while the baseline methods fail to reach -5.0 in 1,400s. In (c), we vary the number of sampling steps per iteration from 5 to 100 (we omit the results of Gibbs-1 since it diverges with less than 100 steps) and report the RMSE after 10,000 iterations. DMALA and DULA outperform GWG consistently and the improvement becomes larger when the number of sampling

steps becomes smaller, demonstrating the fast mixing of our discrete Langevin proposal.

7.3.2. DEEP EBMS

We train deep EBMs where E_θ is a ResNet (He et al., 2016) with Persistent Contrastive Divergence (Tieleman, 2008; Tieleman & Hinton, 2009) and a replay buffer (Du & Mor-datch, 2019) following Grathwohl et al. (2021). We run DMALA and DULA for 40 steps per iteration. After training, we adopt Annealed Importance Sampling (Neal, 2001) to estimate the likelihood. The results for GWG and Gibbs are taken from Grathwohl et al. (2021), and for VAE are taken from Tomczak & Welling (2018).

Results In Table 1, we see that DMALA yields the highest log-likelihood among all methods and its generated images in Figure 9 in Appendix I.4 are very close to the true images. DULA runs the same number of steps as DMALA and GWG but only with half of the cost. We hypothesize that running DULA for more steps or with an adaptive stepsize schedule (Song & Ermon, 2019) will improve its performance.

7.4. Binary Bayesian Neural Networks

Bayesian neural networks have been shown to provide strong predictions and uncertainty estimation in deep learning (Hernández-Lobato & Adams, 2015; Zhang et al., 2020; Liu et al., 2021a). In the meanwhile, binary neural networks (Courbariaux et al., 2016; Rastegari et al., 2016; Liu et al., 2021b), i.e. the weight is in $\{-1, 1\}$, accelerate the learning and significantly reduce computational and memory costs. To combine the benefits of both worlds, we consider training a binary Bayesian neural network with discrete sampling. We conduct regression on four UCI datasets (Dua &

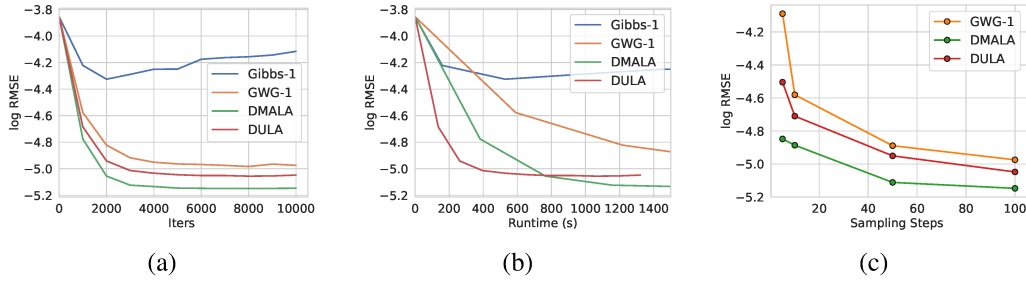


Figure 5. Learning Ising models. DULA and DMALA find W in a shorter time than baselines, and the improvement gets larger with smaller number of sampling steps.

Model	Methods	Self-BLEU ($\downarrow$)	Unique n -grams (%) ($\uparrow$)						Corpus BLEU ($\uparrow$)
			Self		WT103		TBC		
			$n = 2$	$n = 3$	$n = 2$	$n = 3$	$n = 2$	$n = 3$	
Bert-Base	Gibbs	86.84	10.98	16.08	18.57	32.21	21.22	33.05	23.82
	GWG	81.97	15.12	21.79	22.76	37.59	24.72	37.98	22.84
	DULA	72.37	23.33	32.88	27.74	45.85	30.02	46.75	21.82
	DMALA	72.59	23.26	32.64	27.99	45.77	30.32	46.49	21.85
Bert-Large	Gibbs	88.78	9.31	13.74	17.78	30.50	20.48	31.23	22.57
	GWG	86.50	11.03	16.13	19.25	33.20	21.42	33.54	23.08
	DULA	77.96	17.97	26.64	23.69	41.30	26.18	42.14	21.28
	DMALA	76.27	19.83	28.48	25.38	42.94	27.87	43.77	21.73

Table 3. Quantative results on text infilling. The reference text for computing the Corpus BLEU is the combination of WT103 and TBC. DULA and DMALA generate sentences with a similar level of quality as Gibbs and GWG (measured by Corpus BLEU) while attaining higher diversity (measured by self-BLEU and Unique n -grams).

Graff, 2017), and the energy function is defined as,

$$U(\theta) = - \sum_{i=1}^N \|f_{\theta}(x_i) - y_i\|^2,$$

where $D = \{x_i, y_i\}_{i=1}^N$ is the training dataset, and f_{θ} represents a two-layer neural network with $\tanh$ activation and 500 hidden neurons. The dimension of the weight d varies on different datasets ranging from 7,500 to 45,000. We report the log-likelihood on the training set together with root mean-square-error (RMSE) on the test set.

Results From Table 2, we observe that DMALA and DULA outperform other methods significantly on all datasets except test RMSE on COMPAS, which we hypothesize is because of overfitting (the training set only has 4,900 data points). These results demonstrate that our methods converge fast for high dimensional distributions, due to the ability to make large moves per iteration, and suggest that our methods are compelling for training low-precision Bayesian neural networks of which the weight is discrete.

7.5. Text Infilling

Text infilling is an important and intriguing task where the goal is to fill in the blanks given the context (Zhu et al., 2019;

Donahue et al., 2020). Prior work has realized it by sampling from a categorical distribution produced by BERT (Devlin et al., 2019; Wang & Cho, 2019). However, there are a huge number of word combinations, which makes sampling from the categorical distribution difficult. Therefore, an efficient sampler is needed to producing high-quality text.

We randomly sample 20 sentences from TBC (Zhu et al., 2015) and WikiText-103 (Merity et al., 2017), mask 25% of the words in the sentence (Zhu et al., 2019; Donahue et al., 2020), and sample 25 sentences from the probability distribution given by BERT. We run all samplers for 50 steps based on two models, BERT-base and BERT-large. As a common practice in non-autoregressive text generation, we select the top-5 sentences with the highest likelihood out of the 25 sentences to avoid low-quality generation (Gu et al., 2018; Zhou et al., 2019). We evaluate the methods from two perspectives, diversity and quality. For diversity, we use self-BLEU (Zhu et al., 2018) and the number of unique n -grams (Wang & Cho, 2019) to measure the difference between the generated sentences. For quality, we measure the BLEU score (Papineni et al., 2002) between the generated texts and the original dataset (TBC+WikiText-103) (Wang & Cho, 2019; Yu et al., 2017). Note that the BLEU score can only approximately represent the quality of the generation since it cannot handle out-of-distribution generations.

Infilling Task: it also marked the first time the dodgers had won six straight opening day games , [MASK] of which he [MASK] .		
Gibbs Results: all of which he started all of which he started all of which he started all of which he started four of which he started	GWG Results: five of which he started five of which he started four of which he started four of which he pitched two of which he started	DMALA Results: two of which he started one of which he started three of which he led all of which he selected six of which he missed
Infilling Task: he had not , after all , [MASK] me the chance but [MASK] abandoned me [MASK] .		
Gibbs Results: given me the chance but had abandoned me instead given me the chance but had abandoned me instead given me the chance but had abandoned me instead given me the chance but had abandoned me completely given me the chance but had abandoned me anyway	GWG Results: given me the chance but had abandoned me instead given me the chance but had abandoned me himself offered me the chance but had abandoned me completely gave me the chance but had abandoned me anyway given me the chance but he abandoned me instead	DMALA Results: shown me the chance but had abandoned me anyway shown me the chance but not abandoned me immediately gives me the chance but also abandoned me perhaps grants me the chance but really abandoned me entirely offered me the chance but yet abandoned me instead

Figure 6. Examples of the generated sentences on text infilling. Blue and Red words are generated where Red indicates repetitive generation. DMALA can generate semantically meaningful sentences with much higher diversity.

We use one-hot vectors to represent categorical variables (see Appendix B for a discussion of DLP with categorical variables in practice).

Results The quantitative and qualitative results are shown in Table 3 and Figure 6. We find that DULA and DMALA can fill in the blanks with similar quality as Gibbs and GWG but with much higher diversity. Due to the nature of languages, there exist strong correlations among words. It is generally difficult to change one word given the others fixed while still fulfilling the context. However, it is likely to have another combination of words, which are all different from the current ones, to satisfy the infilling. Because of this, the ability to update all coordinates in one step makes our methods especially suitable for this task, as reflected in the evaluation metrics and generated sentences.

8. Conclusion

We propose a Langevin-like proposal for efficiently sampling from complex high-dimensional discrete distributions. Our method, discrete Langevin proposal (DLP), is able to explore discrete structure effectively based on the gradient information. For different usage scenarios, we have developed several variants with DLP, including unadjusted, Metropolis-adjusted, stochastic, and preconditioned versions. We prove the asymptotic convergence of DLP without the MH step under log-quadratic and general distributions. Empirical results on many different problems demonstrate the superiority of our method over baselines in general settings.

While the Langevin algorithm has achieved great success in continuous spaces, there has always lacked a counterpart of such simple, effective and general-purpose samplers in discrete spaces. We hope our method sheds light on building practical and accurate samplers for discrete distributions.

Acknowledgements

We would like to thank Yingzhen Li and the anonymous reviewers for their thoughtful comments on the manuscript. This research is supported by CAREER-1846421, SenSE2037267, EAGER-2041327, Office of Navy Research, and NSF AI Institute for Foundations of Machine Learning (IFML).

References

- Brooks, S., Gelman, A., Jones, G., and Meng, X.-L. *Handbook of markov chain monte carlo*. CRC press, 2011.
- Buza, K. Feedback prediction for blogs. In *Data analysis, machine learning and knowledge discovery*, pp. 145–152. Springer, 2014.
- Cho, G. E. and Meyer, C. D. Comparison of perturbation bounds for the stationary distribution of a markov chain. *Linear Algebra and its Applications*, 335(1-3):137–150, 2001.
- Courbariaux, M., Hubara, I., Soudry, D., El-Yaniv, R., and Bengio, Y. Binarized neural networks: Training deep neural networks with weights and activations constrained to+ 1 or-1. *arXiv preprint arXiv:1602.02830*, 2016.
- Devlin, J., Chang, M.-W., Lee, K., and Toutanova, K. Bert: Pre-training of deep bidirectional transformers for language understanding. *Association for Computational Linguistics*, 2019.
- Donahue, C., Lee, M., and Liang, P. Enabling language models to fill in the blanks. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pp. 2492–2501, 2020.
- Du, Y. and Mordatch, I. Implicit generation and generalization in energy-based models. *Advances in Neural Information Processing Systems*, 2019.

- Dua, D. and Graff, C. UCI machine learning repository, 2017. URL <http://archive.ics.uci.edu/ml>.
- Duane, S., Kennedy, A. D., Pendleton, B. J., and Roweth, D. Hybrid monte carlo. *Physics letters B*, 195(2):216–222, 1987.
- Durmus, A. and Moulines, E. High-dimensional bayesian inference via the unadjusted langevin algorithm. *Bernoulli*, 25(4A):2854–2882, 2019.
- Grathwohl, W., Wang, K.-C., Jacobsen, J.-H., Duvenaud, D., Norouzi, M., and Swersky, K. Your classifier is secretly an energy based model and you should treat it like one. In *International Conference on Learning Representations*, 2019.
- Grathwohl, W., Swersky, K., Hashemi, M., Duvenaud, D., and Maddison, C. J. Oops i took a gradient: Scalable sampling for discrete distributions. *International Conference on Machine Learning*, 2021.
- Grenander, U. and Miller, M. I. Representations of knowledge in complex systems. *Journal of the Royal Statistical Society: Series B (Methodological)*, 56(4):549–581, 1994.
- Gu, J., Bradbury, J., Xiong, C., Li, V. O., and Socher, R. Non-autoregressive neural machine translation. In *International Conference on Learning Representations*, 2018.
- Han, J., Ding, F., Liu, X., Torresani, L., Peng, J., and Liu, Q. Stein variational inference for discrete distributions. In *International Conference on Artificial Intelligence and Statistics*, pp. 4563–4572. PMLR, 2020.
- Hastings, W. K. Monte carlo sampling methods using markov chains and their applications. 1970.
- He, K., Zhang, X., Ren, S., and Sun, J. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 770–778, 2016.
- Hernández-Lobato, J. M. and Adams, R. Probabilistic back-propagation for scalable learning of bayesian neural networks. In *International conference on machine learning*, pp. 1861–1869. PMLR, 2015.
- Hinton, G. E. Training products of experts by minimizing contrastive divergence. *Neural computation*, 14(8):1771–1800, 2002.
- J. Angwin, J. Larson, S. M. and Kirchner, L. Machine bias: There’s software used across the country to predict future criminals. and it’s biased against blacks. *ProPublica*, 2016.
- Jaini, P., Nielsen, D., and Welling, M. Sampling in combinatorial spaces with survae flow augmented mcmc. In *International Conference on Artificial Intelligence and Statistics*, pp. 3349–3357. PMLR, 2021.
- Kingma, D. P. and Ba, J. Adam: A method for stochastic optimization. *International Conference for Learning Representations*, 2015.
- LeCun, Y., Chopra, S., Hadsell, R., Ranzato, M., and Huang, F. A tutorial on energy-based learning. *Predicting structured data*, 1(0), 2006.
- Levin, D. A. and Peres, Y. *Markov chains and mixing times*, volume 107. American Mathematical Soc., 2017.
- Li, C., Chen, C., Carlson, D., and Carin, L. Preconditioned stochastic gradient langevin dynamics for deep neural networks. In *Thirtieth AAAI Conference on Artificial Intelligence*, 2016.
- Liu, X., Tong, X., and Liu, Q. Sampling with trustworthy constraints: A variational gradient framework. *Advances in Neural Information Processing Systems*, 34, 2021a.
- Liu, X., Ye, M., Zhou, D., and Liu, Q. Post-training quantization with multiple points: Mixed precision without mixed precision. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 35, pp. 8697–8705, 2021b.
- Merity, S., Xiong, C., Bradbury, J., and Socher, R. Pointer sentinel mixture models. *International conference on machine learning*, 2017.
- Metropolis, N., Rosenbluth, A. W., Rosenbluth, M. N., Teller, A. H., and Teller, E. Equation of state calculations by fast computing machines. *The journal of chemical physics*, 21(6):1087–1092, 1953.
- Neal, R. M. Annealed importance sampling. *Statistics and computing*, 11(2):125–139, 2001.
- Neal, R. M. et al. Mcmc using hamiltonian dynamics. *Handbook of markov chain monte carlo*, 2(11):2, 2011.
- Nishimura, A., Dunson, D. B., and Lu, J. Discontinuous hamiltonian monte carlo for discrete parameters and discontinuous likelihoods. *Biometrika*, 107(2):365–380, 2020.
- Pakman, A. and Paninski, L. Auxiliary-variable exact hamiltonian monte carlo samplers for binary distributions. *Advances in Neural Information Processing Systems*, 2013.
- Papineni, K., Roukos, S., Ward, T., and Zhu, W.-J. Bleu: a method for automatic evaluation of machine translation. In *Proceedings of the 40th annual meeting of the Association for Computational Linguistics*, pp. 311–318, 2002.

- Peters, J. W. and Welling, M. Probabilistic binary neural networks. *arXiv preprint arXiv:1809.03368*, 2018.
- Power, S. and Goldman, J. V. Accelerated sampling on discrete spaces with non-reversible markov processes. *arXiv preprint arXiv:1912.04681*, 2019.
- Ramachandran, P., Zoph, B., and Le, Q. V. Searching for activation functions. *arXiv preprint arXiv:1710.05941*, 2017.
- Rastegari, M., Ordonez, V., Redmon, J., and Farhadi, A. Xnor-net: Imagenet classification using binary convolutional neural networks. In *European conference on computer vision*, pp. 525–542. Springer, 2016.
- Roberts, G. O. and Stramer, O. Langevin diffusions and metropolis-hastings algorithms. *Methodology and computing in applied probability*, 4(4):337–357, 2002.
- Roberts, G. O. and Tweedie, R. L. Exponential convergence of langevin distributions and their discrete approximations. *Bernoulli*, pp. 341–363, 1996.
- Sansone, E. Lsb: Local self-balancing mcmc in discrete spaces. *arXiv preprint arXiv:2109.03867*, 2021.
- Schweitzer, P. J. Perturbation theory and finite markov chains. *Journal of Applied Probability*, 5(2):401–413, 1968.
- Song, Y. and Ermon, S. Generative modeling by estimating gradients of the data distribution. *Advances in Neural Information Processing Systems*, 2019.
- Sun, H., Dai, H., Xia, W., and Ramamurthy, A. Path auxiliary proposal for mcmc in discrete space. In *International Conference on Learning Representations*, 2022.
- Tieleman, T. Training restricted boltzmann machines using approximations to the likelihood gradient. In *Proceedings of the 25th international conference on Machine learning*, pp. 1064–1071, 2008.
- Tieleman, T. and Hinton, G. Using fast weights to improve persistent contrastive divergence. In *Proceedings of the 26th annual international conference on machine learning*, pp. 1033–1040, 2009.
- Titsias, M. K. and Yau, C. The hamming ball sampler. *Journal of the American Statistical Association*, 112(520): 1598–1611, 2017.
- Tomczak, J. and Welling, M. Vae with a vampprior. In *International Conference on Artificial Intelligence and Statistics*, pp. 1214–1223. PMLR, 2018.
- Wang, A. and Cho, K. Bert has a mouth, and it must speak: Bert as a markov random field language model. *arXiv preprint arXiv:1902.04094*, 2019.
- Wang, J., Huda, A., Lunyak, V. V., and Jordan, I. K. A gibbs sampling strategy applied to the mapping of ambiguous short-sequence tags. *Bioinformatics*, 26(20):2501–2508, 2010.
- Welling, M. and Teh, Y. W. Bayesian learning via stochastic gradient langevin dynamics. In *Proceedings of the 28th international conference on machine learning*, pp. 681–688. Citeseer, 2011.
- Yu, L., Zhang, W., Wang, J., and Yu, Y. Seqgan: Sequence generative adversarial nets with policy gradient. In *Proceedings of the AAAI conference on artificial intelligence*, volume 31, 2017.
- Zanella, G. Informed proposals for local mcmc in discrete spaces. *Journal of the American Statistical Association*, 115(530):852–865, 2020.
- Zhang, D., Malkin, N., Liu, Z., Volokhova, A., Courville, A., and Bengio, Y. Generative flow networks for discrete probabilistic modeling. *arXiv preprint arXiv:2202.01361*, 2022.
- Zhang, R., Li, C., Zhang, J., Chen, C., and Wilson, A. G. Cyclical stochastic gradient mcmc for bayesian deep learning. In *International Conference on Learning Representations*, 2020.
- Zhou, C., Gu, J., and Neubig, G. Understanding knowledge distillation in non-autoregressive machine translation. In *International Conference on Learning Representations*, 2019.
- Zhou, G. Mixed hamiltonian monte carlo for mixed discrete and continuous variables. *Advances in Neural Information Processing Systems*, 33:17094–17104, 2020.
- Zhu, W., Hu, Z., and Xing, E. Text infilling. *arXiv preprint arXiv:1901.00158*, 2019.
- Zhu, Y., Kiros, R., Zemel, R., Salakhutdinov, R., Urtasun, R., Torralba, A., and Fidler, S. Aligning books and movies: Towards story-like visual explanations by watching movies and reading books. In *Proceedings of the IEEE international conference on computer vision*, pp. 19–27, 2015.
- Zhu, Y., Lu, S., Zheng, L., Guo, J., Zhang, W., Wang, J., and Yu, Y. Taxygen: A benchmarking platform for text generation models. In *The 41st International ACM SIGIR Conference on Research & Development in Information Retrieval*, pp. 1097–1100, 2018.