
Stochastic Quantum Sampling for Non-Logconcave Distributions and Estimating Partition Functions

Guneykan Ozgul¹ Xiantao Li² Mehrdad Mahdavi¹ Chunhao Wang¹

Abstract

We present quantum algorithms for sampling from possibly non-logconcave probability distributions expressed as $\pi(x) \propto \exp(-\beta f(x))$ as well as quantum algorithms for estimating the partition function for such distributions. We also incorporate a stochastic gradient oracle that implements the quantum walk operators inexactly by only using mini-batch gradients when f can be written as a finite sum. One challenge of quantizing the resulting Markov chains is that they do not satisfy the detailed balance condition in general. Consequently, the mixing time of the algorithm cannot be expressed in terms of the spectral gap of the transition density matrix, making the quantum algorithms nontrivial to analyze. We overcame these challenges by first building a reference reversible Markov chain that converges to the target distribution, then controlling the discrepancy between our algorithm’s output and the target distribution by using the reference Markov chain as a bridge to establish the total complexity. Our quantum algorithms exhibit polynomial speedups in terms of dimension or precision dependencies when compared to best-known classical algorithms under similar assumptions.

1. Introduction

Many problems in statistics, physics, finance, machine learning, optimization, and molecular dynamics involve sampling from a distribution with a density proportional to $e^{-\beta f(x)}$, known as Gibbs-Boltzmann distribution. For instance, techniques for sampling from such a distribution play a central role in statistical mechanics in probing equilibrium states

of physical systems, understanding phase transition, and estimating thermodynamic properties (Chandler, 1987). In machine learning, sampling from these distributions aids in exploring the posterior distribution in Bayesian learning (Neal et al., 2011; Ahn et al., 2012; Cheng et al., 2018), enabling parameter estimation, uncertainty quantification, and model comparison (Murphy, 2022). In convex geometry, an effective sampling strategy is central to estimating the volume of a convex body that can be applied to problems in statistics, theoretical computer science, and operations research (Vempala, 2007).

A well-known classical method for Gibbs sampling is Markov Chain Monte Carlo (MCMC) method, where a Markov chain with desired stationary density is constructed. Then, the samples can be generated by running the Markov chain for a sufficiently long time. One such Markov chain can be obtained through careful discretization of Langevin diffusion equation and this technique inspired a large family of gradient-based sampling algorithms.

Langevin diffusion is a continuous stochastic differential equation that converges to the desired Gibbs distribution. A simple, yet common, discretization of Langevin diffusion is Euler-Maruyama method with sufficiently small step size $\eta > 0$, resulting in an MCMC algorithm (See Appendix B.3 for more details). However, due to finite-sized discretization, the Markov chain is asymptotically biased. That is, it only converges to the neighborhood of the desired Gibbs distribution, prohibiting one from using large step sizes because of this discrepancy. To overcome this bias, one can adjust the Markov chain by introducing Metropolis-Hastings filter, which is used as a conditional rejection to guarantee that the chain is time-reversible (See Appendix B.1 for the definition of reversibility) and it converges to the desired distribution. This algorithm is sometimes referred to as Metropolis-adjusted Langevin algorithm (MALA), and the algorithm without the rejection step is conventionally called unadjusted Langevin algorithm (ULA). We refer the reader to Appendix B for more details about these algorithms.

In the past decade, notable progress has been witnessed in the theoretical development of quantum algorithms for various machine learning and optimization problems. It is natural to expect that quantum computers also provide

¹Department of Computer Science and Engineering, Pennsylvania State University ²Department of Mathematics, Pennsylvania State University. Correspondence to: Guneykan Ozgul <gmo5119@psu.edu>, Chunhao Wang <cwang@psu.edu>.

provable speedups for general sampling problems. If we could prepare a quantum state whose amplitudes correspond to some desired distribution, then measuring this state yields a random sample from this probability distribution. Unfortunately, quantum speedups in such sampling models probably do not hold in general as this will imply $\text{SZK} \subseteq \text{BQP}$ (Aharonov & Ta-Shma, 2003). While the hardness barrier exists for a quantum speedup for general sampling problems, in some special cases, it has been shown that quantum algorithms can achieve polynomial speedups over classical algorithms. Such examples include quantum algorithms for uniform sampling on a 2D lattice (Richter, 2007), for estimating partition functions (Wocjan & Abeyesinghe, 2008; Wocjan et al., 2009; Montanaro, 2015; Harrow & Wei, 2020; Arunachalam et al., 2021; Cornelissen & Hamoudi, 2023), and for estimating volumes of convex bodies (Chakrabarti et al., 2023).

Recently, a quantum MALA algorithm based on quantum simulated annealing is introduced (Childs et al., 2022), which leverages the fact that a coherent quantum state corresponding to desired logconcave distribution can be prepared using fewer number of calls to gradient and evaluation oracle than the classical counterparts. Inspired by this, an interesting question arises: *Can we attain quantum speedups for more general distributions, such as non-logconcave distributions?* Moreover, one intriguing open question posed in (Childs et al., 2022) was the possibility of speeding up unadjusted Langevin algorithm using similar techniques. The **main challenge** for analyzing quantum version of ULA (or stochastic ULA) is that the transition density does not satisfy the detailed balance condition due to lack of the Metropolis-Hastings filter. Hence the Markov chain is not time reversible which is the main assumption for almost all quantum walk based algorithms (Szegedy, 2004; Wocjan & Abeyesinghe, 2008; Magniez et al., 2011; Apers & Sarlette, 2019). The current quantum walk frameworks leverage the fact that a symmetric discriminant matrix $D_{xy} = \sqrt{P_{xy}P_{yx}}$ can be related to the spectrum of the classical transition matrix P . Then the eigenstate of D with unique singular value 1 encodes the coherent quantum state whose amplitudes are the desired Gibbs density. Then, by extracting this eigenstate using a quantum computer can prepare the Gibbs state quadratically faster than the classical computers in spectral gap parameter. However, for non-reversible Markov chains the spectral connection between D and P are not straightforward. In fact, the discriminant matrix for non-reversible chains might have zero singular value gap (Magniez et al., 2011), which breaks down the entire quantum algorithm. We also touch upon this technical difficulty of quantizing non-reversible Markov chains in Section 4 in more detail.

One particular reason to analyze ULA and stochastic ULA is that implementing one step of MALA requires N function and gradient evaluations when f can be decomposed into

a sum of N terms, whereas ULA only uses N gradient evaluations as it does not do any adjustment. We prove that stochastic ULA only needs $O(d)$ gradient evaluations to converge which is a significant improvement especially when $d \ll N$. Therefore, we believe each algorithm is suited to specific use cases, depending on the size of the data set, problem dimension, and hardness of function and gradient evaluations.

Main Contributions

- We analyzed the mixing time of quantum MALA algorithm (Theorem 4.1) for non-logconcave distributions, extending the work done in (Childs et al., 2022). The main challenge in analyzing quantum MALA for non-logconcave distributions is to characterize the phase gap of the quantum walk and to show the existence of a quantum annealing schedule that guarantees a large overlap between successive distributions since the target distribution does not satisfy the concentration inequalities as in log-concave case. By using the conductance analysis done in (Zou et al., 2021), we characterized the phase gap. Next, we showed that by using isoperimetric inequalities, the length of the annealing schedule is $\tilde{O}(\sqrt{d})$ ¹ similar to non-logconcave case (Section 3).
- We analyzed quantum ULA algorithm (Theorem 4.2) using a novel perturbation analysis with respect to quantum MALA to show, for the first time, that quantum computers can provide speedups even for non-reversible Markov chains. Since quantum MALA is time-reversible (as it satisfies Equation (34)) and asymptotically unbiased, it converges to target distribution, allowing us to express our algorithms' error with respect to Gibbs distribution. In the construction of our algorithms, we use standard quantum simulated annealing techniques as in (Childs et al., 2022) while the underlying Markov chain is non-reversible. Although perturbation techniques have been used in classical analysis of Markov chains (Zou et al., 2021; Raginsky et al., 2017; Xu et al., 2018), these results cannot be transferred to quantum setting as the quantum Monte Carlo algorithms are fundamentally different. That is, while classical algorithms run in an iterative fashion to generate candidate samples, quantum algorithms use linear algebraic techniques to rotate the input state towards the eigenvector in the invariant subspace. We believe this technique can be useful for the analysis of other non-reversible Markov chains as an independent tool.
- We further incorporated stochastic gradient oracle to make the implementation of quantum walk efficient and provided the mixing time of our stochastic quantum sampling algorithm in Theorem 4.3. In addition to the error due to

¹Throughout this paper, we use the notation $\tilde{O}(\cdot)$ to hide the poly-logarithmic dependencies on β , ϵ , d , ρ , and CLSI .

the lack of Metropolis-Hastings filter, the stochastic algorithm introduces additional errors because of the noisy gradients. We use concentration techniques to show that even with noisy unitaries, the quantum algorithm gives the correct distribution with high probability.

- We combined our sampling algorithms with recently developed efficient quantum product estimator (Cornelissen & Hamoudi, 2023) and proposed algorithms for computing the partition function for non-logconcave distributions in Section 5. Our algorithm might have other application areas such as computing volumes for non-convex bodies although this would require additional assumptions and techniques to express the volume estimation as an approximate counting problem where one can use partition function estimation to solve.

Problem formulation In this paper, we focus on designing and analyzing quantum algorithms for sampling from the Gibbs density $\propto e^{-\beta f(x)}$ where $f(x)$ is not necessarily convex. An important scenario in machine learning is when $f(x)$ admits a decomposition,

$$f(x) = \frac{1}{N} \sum_{k=1}^N f_k(x), \quad (1)$$

where $N \gg 1$ is large. One typical example is where x comes from the model parameter, and $f(x)$ is the empirical loss defined on a large data set. Clearly, this will cause a significant slowdown of quantum MALA algorithm due to function and gradient evaluations when f is given in this finite sum form. We further make the following assumptions on f . These assumptions are realistic as they are satisfied in many applications and are widely assumed in the non-logconcave sampling literature (Raginsky et al., 2017; Zou et al., 2021).

Assumption 1.1 (Smoothness). There exists a positive constant L such that for any $x, y \in \mathbb{R}^d$ and all functions $f_k(x), k \in [N]$, it holds that

$$\|\nabla f_k(x) - \nabla f_k(y)\| \leq L\|x - y\|. \quad (2)$$

Assumption 1.2 (Dissipativeness). There are absolute constants $m > 0$ and $b \geq 0$ such that

$$\langle \nabla f(x), x \rangle \geq m\|x\|^2 - b. \quad (3)$$

The first assumption ensures that small changes in the input parameters result in bounded changes in gradients whereas the second one implies that $f(x)$ grows like a quadratic form outside a ball.

Remark 1.3. Classical works used in Table 1 has slightly different assumptions than Assumptions 1.1 and 1.2. In the work of Ma et al. (2019b), the function f is assumed to be

strongly convex outside a sufficiently large ball with radius. This shares the same intuition with the dissipative condition, i.e., sufficiently fast growth in the far field. Moreover it implies the dissipative condition. The isoperimetry condition in Vempala & Wibisono (2019) relies on log-Sobolev inequality, which can be proved from the dissipative condition together with the Lipschitz condition (See proposition 3.2 in Raginsky et al. (2017)). In fact, this is why we do not add log-Sobolev equality as an additional assumption.

Next, we give the following definitions that are commonly used in the analysis of non-logconcave sampling.

Definition 1.4 (Cheeger Constant). Let ν be a probability measure on Ω . Then ν satisfies the isoperimetric inequality with Cheeger constant ρ if for any $A \subseteq \Omega$, it holds that

$$\liminf_{h \rightarrow 0^+} \frac{\nu(A_h) - \nu(A)}{h} \geq \rho \min\{\nu(A), 1 - \nu(A)\}, \quad (4)$$

where $A_h = \{x \in \Omega : \exists y \in A, \|x - y\| \leq h\}$.

Definition 1.5 (Log-Sobolev Inequality). Let ν be a probability measure on Ω . We say that ν satisfies log-Sobolev inequality with constant c_{LSI} if for any smooth function g on $\mathbb{R}^d$, satisfying $\int_{\Omega} g(x)\nu(x) dx = 1$, it holds that

$$\int_{\Omega} g(x) \log(g(x)) \nu(x) dx \leq \frac{1}{2c_{\text{LSI}}} \int_{\Omega} \frac{\|\nabla g(x)\|^2}{g(x)} \nu(x) dx. \quad (5)$$

While the Cheeger constant measures the bottleneck of a space, the log-Sobolev constant resembles the PL (Polyak-Łojasiewicz) constant in optimization space. That is, it quantifies the distance between an iterate to optimum in terms of the norm of the gradient.

Oracle model We assume that we have the access to the following quantum oracles to implement our algorithm. These oracles are virtually classical oracles while empowering superposition access. We first define the *full gradient oracle* for f as follows:

$$\mathcal{O}_{\nabla f} |x\rangle |0\rangle = |x\rangle |\nabla f(x)\rangle. \quad (6)$$

Similarly, we define a *stochastic gradient oracle*,

$$\mathcal{O}_{\tilde{\nabla} f} |x\rangle |0\rangle = |x\rangle |\tilde{\nabla} f(x)\rangle. \quad (7)$$

where $\tilde{\nabla} f(x) = \frac{1}{B} \sum_{k \in S} \nabla f_k(x)$ where S is a subset of size B data samples chosen randomly without replacement. Note that $\mathcal{O}_{\tilde{\nabla} f}$ possibly outputs a different state for the same input state depending on the internal random batch. Finally, the *evaluation oracle* is defined by

$$\mathcal{O}_f |x\rangle |0\rangle = |x\rangle |f(x)\rangle. \quad (8)$$

We note that although we quantified the complexity of our algorithm in terms of the number of calls to these oracles, the evaluation oracle and full gradient oracle are slower than the stochastic gradient oracle for finite sum form due to evaluation of N terms. One of our contributions is to use a stochastic gradient oracle to make quantum walk implementation more efficient.

We also emphasize that our gradient oracles, including those for stochastic gradients, operate classically with superposition access. Since classical circuits can be simulated by quantum circuits with a constant overhead, implementation cost of these oracles is on par with the classical oracles. Therefore, any speedups with respect to the number of calls to these oracles are not suppressed by their implementation cost. Furthermore, our quantum algorithms are robust to a small error in these oracles. This is because a small error in gradient will introduce a perturbation of the Markov chain and the resulting quantum walk. The analysis in this paper, which is precisely based on quantifying the difference between two Markov chains (a time-irreversible and a time-reversible chain), can quantify how the gradient error can propagate in the algorithm. In fact, our analysis for quantum ULA with stochastic gradients show that the algorithm works for even noisy gradients.

Notation The notation $\|\cdot\|$ denotes the spectral norm for operators and the ℓ^2 norm for quantum states. For Markov chains, we use the notation $P(x, \cdot)$ to denote the transition probability distribution for point $x \in \Omega$, whereas we use $P(x, y)$ or p_{xy} to denote the probability of transitioning from point x to y . For a distribution $p(x)$ and a function $q(x)$, the notation $p(x) \propto q(x)$ means p is proportional to q up to a normalization factor. The ket notation $|\nu\rangle$ is sometimes referred to the coherent quantum state corresponding to probability distribution ν and is not explicitly stated when it is clear from the context. For distributions, $\|\cdot\|_{TV}$ denotes the total variation distance, $\|\cdot\|_H$ Hellinger distance and finally $W_2(\cdot, \cdot)$ denotes the Wasserstein-2 distance. The total variation distance between two probability distributions P and Q on Ω is defined as:

$$\|P - Q\|_{TV} := \sup_{A \subseteq \Omega} |P(A) - Q(A)|, \quad (9)$$

and the Hellinger distance is defined as,

$$\|P - Q\|_H := \left(\frac{1}{2} \int_{x \in \Omega} \left(\sqrt{P(dx)} - \sqrt{Q(dx)} \right)^2 \right)^{1/2}, \quad (10)$$

and finally, Wasserstein-2 distance is defined as,

$$W_2(P, Q) := \left(\inf_{z \in \Gamma(P, Q)} \int_{\mathbb{R}^d \times \mathbb{R}^d} \|x - y\|^2 dz(x, y) \right)^{1/2} \quad (11)$$

where Γ is the set of all couplings between P and Q .

2. Related Work

Extensive research has been conducted to understand the non-asymptotic dynamics of Langevin based algorithms for both log-concave and non-logconcave densities under various settings. This section reviews a selection of significant works to provide context for our study, given the extensive literature available.

For log-concave distributions, a significant body of research has been conducted to understand the dynamics of the Langevin Monte Carlo (LMC) based algorithms (Bubeck et al., 2015; Dalalyan, 2017a;b; Durmus et al., 2019; Li et al., 2022). Sampling from non-logconcave distributions under various assumptions have also been analyzed broadly (Lee et al., 2018; Vempala & Wibisono, 2019; Ma et al., 2019b; Xu et al., 2018). The convergence of LMC under the condition that the target density satisfies isoperimetry condition is shown by Vempala & Wibisono (2019). Although the gradient descent methods are known to be superior to sampling-based optimization in convex cases, Ma et al. (2019b) showed that sampling-based methods could provide speedups over local optimization methods in non-convex setting which motivates us to explore the quantum algorithms for non-logconcave densities. The stochastic extension of the algorithm (SGLD) in non-logconcave setting has been investigated recently in several works. The hitting time of the stochastic Langevin dynamics to a neighborhood of the minima is analyzed in Zhang et al. (2017) and they showed that SGLD can escape suboptimal local minima that only exist in the empirical risk function. More notably Zou et al. (2021) and Xu et al. (2018) analyzed the mixing time of SGLD to the stationary distribution in total variation distance using similar techniques to ours. They used perturbation analysis to show that output of SGLD is closed to a reversible chain. Unfortunately, their result does not transfer to the quantum setting due to fundamental difference between quantum and classical Monte Carlo algorithms.

We also note the classical sampling algorithms that uses more sophisticated techniques to improve the mixing time in terms of various distances. One such popular technique is called Hamiltonian Monte Carlo method (Brooks et al., 2011) which uses the momentum and leapfrog integrator to reduce the error of discretization which improves the sampling time. Based on underdamped Langevin Monte Carlo algorithm, (Shen & Lee, 2019) proposed randomized midpoint method to sample from log-concave distributions with better dependencies compared to unadjusted Langevin algorithm. Furthermore, (Fan et al., 2023) used proximal sampling algorithm to improve the dimension dependence to $d^{1/2}$ under the log-Sobolev inequality. However, their assumption is ℓ^1 smoothness which differs from our Assumption 1.1.

We only compare our results to the classical ones that use

Table 1. The comparison of our sampling algorithm to classical results with similar assumptions. Here we focus the dependencies on d and ϵ (See the discussion for other parameters Appendix F).

Algorithm	Query Complexity	Oracle	Assumptions
ULA (Ma et al., 2019b)	$\tilde{O}(d/\epsilon^2)$	Full Gradient	Local non-convex
MALA (Ma et al., 2019b)	$\tilde{O}(d^2)$	Full Gradient and Evaluation	Local non-convex
ULA (Vempala & Wibisono, 2019)	$\tilde{O}(d/\epsilon^2)$	Full Gradient	Isoperimetry
SGLD (Zou et al., 2021)	$\tilde{O}(d^4/\epsilon^2)$	Stochastic Gradient	Dissipative Gradients
Quantum MALA (Theorem 4.1)	$\tilde{O}(d)$	Full Gradient and Evaluation	Dissipative Gradients
Quantum ULA (Theorem 4.2)	$\tilde{O}(d^{3/2}/\epsilon)$	Full Gradient	Dissipative Gradients
Stochastic Quantum ULA (Theorem 4.3)	$\tilde{O}(d^{5/2}/\epsilon^2)$	Stochastic Gradient	Dissipative Gradients

the same or very similar assumptions and only claim our polynomial speedups with respect to these results. The comparison is summarized in Table 1.

The quantum walk operators used in this paper is developed in (Szegedy, 2004). Although it is not easy to speed up the mixing time of a general random walks, (Wocjan & Abeyesinghe, 2008) showed that for slowly varying Markov chains, it is possible to achieve quadratic speed up if the stationary distributions of successive Markov chains have large overlap and the initial distribution could be prepared efficiently. This technique has recently been used in quantum optimization problems such as estimating the partition function (Montanaro, 2015), or volume estimation of convex bodies (Chakrabarti et al., 2023). The speedup of non-reversible Markov chains has been discussed by Magniez et al. (2011) in the context of quantum search. However, their construction implements time reversal of the Markov chain $p_{yx}^* = p_{xy}\pi(x)/\pi(y)$ which requires N function evaluations in our case. Therefore, this construction would lead to an algorithm somewhat similar to quantum MALA algorithm.

3. Annealing Schedule for Non-Logconcave Distributions

Wocjan & Abeyesinghe (2008) showed that it is possible to speedup classical Markov chains by using slowly changing Markov chains. Their constructions is first to define a series of Markov chains with stationary distributions $\pi_1, \dots, \pi_M$. Then for each Markov chain, they implement the quantum walk to iteratively drive the initial state $|\pi_1\rangle$ to final state $|\pi_M\rangle$ using amplitude amplification. If the overlap $|\langle\pi_i|\pi_{i+1}\rangle| \geq \Omega(1)$ for all $i \in [1, M-1]$, then the cost of the algorithm becomes M times the cost of implementing quantum walk for each Markov chain.

To implement slowly varying Markov chains for non-logconcave distributions, we prove the following lemma to construct the annealing schedule. It shows that there exists an annealing schedule of length $\tilde{O}(\sqrt{d})$ such that

the adjacent quantum states have large overlap. Our construction is similar to Ge et al. (2020) in that we start from quantum Gaussian state and slowly decrease the Gaussian component of the distribution so that final state is very close to target Gibbs state.

Lemma 3.1 (Quantum Annealing). *Under Assumptions 1.1 and 1.2, there exists a series of quantum states $|\mu_0\rangle, |\mu_1\rangle, \dots, |\mu_M\rangle$ satisfying the following properties:*

1. *There exists an efficient quantum algorithm to prepare initial state $|\mu_0\rangle$ without using any function queries.*
2. *For all $i \in \{0, \dots, M-1\}$, $|\mu_i\rangle$ and $|\mu_{i+1}\rangle$ has at least constant overlap, i.e.,*

$$|\langle\mu_i|\mu_{i+1}\rangle| \geq \Omega(1). \quad (12)$$

3. *The final state $|\mu_M\rangle$ has at least constant overlap with the target Gibbs state $|\pi\rangle$,*

$$|\langle\mu_M|\pi\rangle| \geq \Omega(1). \quad (13)$$

4. *The number of quantum states $M \leq \tilde{O}(c_{\text{LSI}}^{-1}d^{1/2})$.*

We give the details of this schedule and the proofs in Appendix C. We believe, our analysis for this annealing schedule for non-logconcave distributions can have other applications. For instance, Li & Zhang (2022) used an annealing schedule to optimize approximately-convex functions and they showed applications for stochastic bandits. In their construction, they assumed that the objective function can be written as a uniform perturbation of a convex function in the entire domain. Our construction can allow design of optimization algorithms for more general non-convex functions as Assumption 1.1 and Assumption 1.2 are not too restrictive.

4. Quantum Algorithms for Sampling

4.1. Tools from quantum computation

Quantum sampling problem is to create the coherent version of the desired probability distribution:

$$|\pi\rangle = \int_{x \in \mathbb{R}^d} dx \sqrt{\pi(x)} |x\rangle. \quad (14)$$

Then a measurement on this state yields the basis $|x'\rangle$ with probability $\pi(x')$. Let $\pi^{1/2}$ denote the diagonal matrix with entries $\sqrt{\pi(x)}$. For a time reversible Markov chain, the discriminant matrix D is defined by,

$$D(P) = \pi^{1/2} P \pi^{-1/2}. \quad (15)$$

This follows from the detailed balance condition as follows:

$$D_{xy} = \sum_{z,l} \pi_{xz}^{1/2} P_{zl} \pi_{ly}^{-1/2} \quad (16)$$

$$= \sqrt{\pi(x)} P_{xy} \sqrt{\pi(y)} \quad (17)$$

$$= \sqrt{P_{xy} P_{yx}}. \quad (18)$$

Therefore, the spectrum of P matches the spectrum of D . Furthermore,

$$D(P) |\pi\rangle = |\pi\rangle. \quad (19)$$

Therefore the state $|\pi\rangle$ is an eigenvector of $D(P)$ with eigenvalue 1. Then eigenstate of D with eigenvalue 1 on the relevant subspace can then be prepared using singular value transformations (Gilyén et al., 2019) or phase estimation (Wocjan et al., 2009). We use the second approach in this paper. For non-reversible Markov chains we cannot write D in the form of Equation (15), therefore the connection between $|\pi\rangle$ and D is broken. Therefore, in general one cannot expect to prepare $|\pi\rangle$ using quantum linear algebra techniques in the similar fashion.

Quantum walk Classical Markov chains can be quantized on a quantum computer using Szegedy's quantum walk operators introduced in (Szegedy, 2004) by constructing a unitary operator on $\mathcal{H} = \mathbb{C}^N \otimes \mathbb{C}^N$. To be able to implement the quantum walk, one needs to implement the following mapping on a quantum computer.

$$|x\rangle |0\rangle \mapsto |\psi_x\rangle \quad (20)$$

where

$$|\psi_x\rangle = \int_y dy \sqrt{p_{xy}} |x\rangle |y\rangle \quad (21)$$

The unitary operator U can be realized as:

$$U := S \left(2 \int_x |\psi_x\rangle \langle \psi_x| - I \right), \quad (22)$$

where $S = \int_{xy} |x\rangle |y\rangle \langle x| \langle y|$ is the swap operator. Next we show how to implement this mapping described above.

Implementing quantum walk operators We describe the implementation of quantum walk operator for stochastic case here. For full gradient case, we just need to replace the oracle to full gradient oracle. We use the stochastic gradient oracle $O_{\tilde{\nabla}f}$ to prepare the following state,

$$|x\rangle |0\rangle \mapsto |x\rangle |\tilde{\nabla}f(x)\rangle. \quad (23)$$

Since the transition density of unadjusted Langevin algorithm is Gaussian, then one step of the walk can be implemented efficiently on a quantum computer using first Box-Muller transformation (Chakrabarti et al., 2023) and applying a shift operation based on the gradient.

$$|x\rangle |0\rangle \mapsto |x\rangle \int_{\mathbb{R}^d} dy \sqrt{p_{xy}} |y\rangle, \quad (24)$$

where

$$p_{xy} = \left(\frac{1}{2\pi} \right)^{d/2} e^{-\frac{1}{2} \|y - \eta \tilde{\nabla}f(x)\|_2^2}. \quad (25)$$

The query complexity of this operation is $O(B)$ due to B gradient evaluations required to implement the oracle in Equation (23). Note that we present the quantum states and operators in the continuous-space representation. The analysis in continuous-space simplifies the analysis, while the implementations are always in a discretized space (as we only have finite bits of precision for real numbers). We refer to (Chakrabarti et al., 2023) for the error analysis caused by the discretization, which is not dominating other errors. For quantum MALA, we update the target register conditionally similar to (Childs et al., 2022).

Implementing reflection operators To implement amplitude amplification, one needs reflection operator around the target state $|\pi\rangle$. This reflection operator can be approximately implemented using phase estimation circuit (Wocjan & Abeyesinghe, 2008) by using $\tilde{O}(1/\sqrt{\gamma})$ calls to controlled U operators where γ is the spectral gap of the transition density matrix. If U is a quantum walk corresponding to quantum MALA algorithm, we can directly amplify the amplitude of $|\pi\rangle$ thanks to reversibility. However, in quantum ULA or stochastic ULA, the reflection operators are implemented with a bias. Consequently, we converge to a state in the neighborhood of $|\pi\rangle$. We quantified this discrepancy and bound the step size for sufficiently small error.

Amplitude amplification Once we have the appropriate reflection operators, we can use fixed point amplitude amplification technique introduced in (Grover, 2005) to drive the initial state to the target state by applying reflection operators iteratively. It is also possible to use quantum Zeno effect to keep the quantum state close to the Gibbs density by using projective measurements after each phase estimation

(Somma et al., 2007; 2008), it results in worse dependency on the schedule length than the amplitude amplification.

A note on the practicality of the quantum sampling As opposed to a family of quantum machine learning algorithms where the data needs to be encoded into a quantum state, our sampling algorithms uses the data by the gradient and evaluation oracles which are simulated classically. Since the unitary evolutions only use constant number of calls to these oracles, our speedups are not suppressed by other hidden costs such as input preparations etc. However, the overall implementation of quantum sampling algorithm still requires a fault-tolerant quantum computer that can implement the phase estimation circuit with high fidelity. It is an open question whether these speedups can be implemented in near term quantum computers.

Next, we present our results for the sampling algorithms.

4.2. Quantum Metropolis Adjusted Langevin Algorithm

The following theorem establishes the query complexity of quantum MALA algorithm. Since it is a time reversible Markov chain, this result is obtained by characterizing the spectral gap of its transition density (Appendix D.1) using conductance analysis. Then the phase gap of the quantum MALA algorithm scales as $1/\eta^{1/2}$ for sufficiently small step size η . Once the phase gap is characterized, the rest of the proof is to combine it with the annealing schedule.

Theorem 4.1 (Quantum MALA). *Let $\pi \propto e^{-\beta f(x)}$ denote a probability distribution with inverse temperature $\beta > 0$ such that $f(x)$ satisfies Assumptions 1.1 and 1.2. Then, there exists a quantum algorithm that outputs a random variable distributed according to μ such that,*

$$\|\mu - \pi\|_{\text{TV}} \leq \epsilon, \quad (26)$$

where $\|\cdot\|_{\text{TV}}$ is the total variation distance, using $\tilde{O}(\beta d \rho^{-1} c_{\text{LSI}}^{-1})$ queries to $\mathcal{O}_{\nabla f}$ and $\mathcal{O}_f$.

4.3. Quantum Unadjusted Langevin Algorithm

Since unadjusted Langevin algorithm is not a reversible chain, we cannot follow the same procedure since there is no direct relation between the spectral gap and the mixing time of the quantum algorithm. We defer the proof of this theorem to Appendix D.2 and give only the proof sketch here. Our quantum algorithm follows the same procedure as in quantum MALA algorithm, however we implement the quantum walk operators using the transition density of ULA algorithm instead of MALA algorithm. Let U^* and U be the quantum walk operator associated with quantum MALA and quantum ULA respectively. The key idea in our proof is to show that for sufficiently small step size η , the operator norm of the difference $\|U^* - U\| \leq \tilde{O}(\eta d)$. Then, using this one step error between quantum walks and

due to the fact that error accumulates at most linearly with K , the total discrepancy between two algorithms becomes $\tilde{O}(\eta d K)$ where K is the total number of calls to U^* in quantum MALA algorithm with the same step size. Finally, we set the step size sufficiently small so that the total error between two algorithms are smaller than ϵ . Since K is proportional to $1/\eta^{1/2}$, this allows us to characterize K .

Theorem 4.2 (Quantum ULA). *Let $\pi \propto e^{-\beta f(x)}$ denote a probability distribution with inverse temperature $\beta > 0$ such that $f(x)$ satisfies Assumptions 1.1 and 1.2. Then, there exists a quantum algorithm that outputs a random variable distributed according to μ such that,*

$$\|\mu - \pi\|_{\text{TV}} \leq \epsilon, \quad (27)$$

where $\|\cdot\|_{\text{TV}}$ is the total variation distance, using $\tilde{O}(\beta d^{3/2} \epsilon^{-1} \rho^{-1} c_{\text{LSI}}^{-1})$ queries to $\mathcal{O}_{\nabla f}$.

4.4. Quantum Unadjusted Langevin Algorithm with Stochastic Gradients

The construction for the stochastic quantum ULA algorithm is similar to quantum ULA. The stochastic ingredient here is realized by replacing full gradient ∇f with a stochastic gradient $g_\ell = \frac{1}{B} \sum_{k \in S_\ell} \nabla f_k$ in implementing the quantum walk operator where S_ℓ is a batch randomly uniformly drawn from the set $\{A \subseteq [N] : |A| = B\}$. The next theorem, proved in Appendix D.3, quantifies the query complexity of stochastic quantum Langevin algorithm with respect to stochastic gradient oracle. The proof of the query complexity is similar to ULA, however, due to noisy gradients, we need to use matrix concentration to show that the quantum walks are close to each other with high probability for sufficiently small step size. The rigorous proof of the following theorem is deferred to Appendix D.3.

Theorem 4.3 (Quantum ULA with stochastic gradient). *Let $\pi \propto e^{-\beta f(x)}$ denote a probability distribution with inverse temperature $\beta > 0$ such that $f(x) = \frac{1}{N} \sum_{k=1}^N f_k(x)$ satisfies Assumptions 1.1 and 1.2. Then, there exists a quantum algorithm that outputs a random variable distributed according to μ such that,*

$$\|\mu - \pi\|_{\text{TV}} \leq \epsilon, \quad (28)$$

where $\|\cdot\|_{\text{TV}}$ is the total variation distance, using $\tilde{O}(\beta^2 d^{3/2} \epsilon^{-2} \rho^{-2} c_{\text{LSI}}^{-1})^2$ queries to $\mathcal{O}_{\nabla f}$ and each $\mathcal{O}_{\nabla f}$ involves $O(d)$ gradient calculations.

Remark 4.4. The Cheeger and log-Sobolev constants depend on the function landscape and might have different dependence on problem dimension depending on the underlying assumption. For instance, assuming that the function

²As each $\mathcal{O}_{\nabla f}$ uses $\tilde{O}(d)$ gradient calculations, the number of total gradient calculations scale as $d^{5/2}$ as shown Table 1.

is locally non-convex gives a dimension independent c_{LSI} , whereas in general they might have exponential dependence. See (Zou et al., 2021) for a more detailed discussion on these constants. We discuss the dependence of the quantum algorithm on these constants in more detail in Appendix F.

Remark 4.5. In a special scenario where an initial quantum state that has at least constant overlap with $|\pi\rangle$ is provided (e.g. a constant warm state), it is possible to obtain an additional speed up in d dependence by saving up to $O(d^{1/2}c_{\text{LSI}}^{-1})$ using a single Markov chain instead of using simulated annealing.

5. Partition Function Estimation

Computing the partition functions in low temperature regime is a challenging problem that has applications in convex geometry (Chakrabarti et al., 2023), linear algebra (Jerrum et al., 2004), and graph theory (Stefankovic et al., 2009). Even though computing the partition function exactly is a hard problem, it can be approximated up to a multiplicative constant using MCMC methods.

In this section, we describe the method and analysis for estimating the partition function for a non-logconcave distribution defined as,

$$Z = \int_{x \in \mathbb{R}^d} e^{-f(x)} dx. \quad (29)$$

The partition function can be estimated using the following telescoping product:

$$Z = Z_1 \prod_{i=1}^M \frac{Z_{i+1}}{Z_i}, \quad (30)$$

where Z_i is the normalizing constant of the distribution $\mu_i \propto \exp\left(-\frac{\|x\|^2}{2\sigma_i^2} + f(x)\right)$ and $Z_{M+1} = Z$ where $\sigma_1 \leq \sigma_2 \leq \dots \leq \sigma_M$ and $\sigma_{M+1} = \infty$. We then approximate $Z_{M+1} = Z_1 \prod_{i=1}^M \frac{Z_{i+1}}{Z_i} = Z_1 \prod_{i=1}^M \mathbb{E}_{\mu_i}[g_i]$, where

$$g_i = \exp\left(\frac{1}{2} \left(\frac{1}{\sigma_i^2} - \frac{1}{\sigma_{i+1}^2} \right) \|x\|^2\right) \quad (31)$$

for $i \in [M]$ where σ_i is defined in proof of Lemma 3.1. To be able estimate this product, we use the technique proposed by Cornelissen & Hamoudi (2023). Their idea is to estimate each expectation in the product using nearly unbiased quantum mean estimation. Since each term in the product can be estimated faster on a quantum computer than the classical counterparts, the overall algorithm both exploits the fast mean estimation and sampling.

Theorem 5.1. Let $Z = \int_x e^{-f(x)} dx$ be the partition with $f(x)$ function satisfying assumptions Assumptions 1.1

and 1.2. Then, there exists quantum algorithms that output an estimate $\tilde{Z}$ such that,

$$(1 - \epsilon)Z \leq \tilde{Z} \leq (1 + \epsilon)Z \quad (32)$$

with probability at least $3/4$ using,

- $\tilde{O}(d^{5/4}\epsilon^{-1}\rho^{-1}c_{\text{LSI}}^{-1})$ queries to $\mathcal{O}_{\nabla f}$ and $\mathcal{O}_f$, or
- $\tilde{O}(d^{7/4}\epsilon^{-2}\rho^{-1}c_{\text{LSI}}^{-1})$ queries to $\mathcal{O}_{\nabla f}$, or
- $\tilde{O}(d^{11/4}\epsilon^{-3}\rho^{-2}c_{\text{LSI}}^{-1})$ queries to $\mathcal{O}_{\tilde{\nabla} f}$.

We defer the rigorous proof of this theorem to Appendix E. Unfortunately, we are not aware of any classical algorithm for computing the partition function under the same assumptions as ours, therefore we are unable to make a solid comparison.

6. Conclusion and Outlook

We have analyzed algorithms for quantum sampling and estimating partition functions for non-logconcave distributions by quantizing popular techniques in classical sampling literature. We believe our techniques and analysis can be useful tools for developing future quantum Monte Carlo algorithms especially based on non-reversible chains. We list the following theoretical open problems for future work.

- Our quantum algorithms utilize the first order sampling methods used in classical literature. It is known that underdamped Langevin algorithm is the accelerated variant of sampling similar to Nesterov’s acceleration in optimization (Ma et al., 2019a). It is an interesting direction to analyze the possible quantum speedups using such sophisticated classical techniques. Analyzing these possible quantum algorithms in terms of other distance metrics such as Wasserstein or KL distance is also another challenge as these distance metrics are not invariant under unitary transformations.
- Langevin Monte Carlo algorithm is obtained by discretization of the continuous stochastic differential equation known as Langevin diffusion. Analyzing the continuous SPDE in quantum domain directly might be another way of getting around reversibility issue and we might obtain more efficient quantum algorithms.
- Fast forwarding of quantum Markov chains to obtain the transient dynamics rather than its stationary density is also interesting direction and Apers & Sarlette (2019) proposed a quantum algorithm to solve this problem for reversible chains. Using similar perturbation analysis can potentially be used to show that non-reversible chains can also be fast-forwarded under special settings faster than classical counterparts.

Acknowledgements

GO was supported by a seed grant from the Institute of Computational and Data Science (ICDS). XL’s research is supported by the National Science Foundation Grants DMS-2111221. CW acknowledges support from National Science Foundation grant CCF-2238766 (CAREER). We thank Yassine Hamoudi for suggesting to use unbiased mean estimation to improve the result on partition function estimation in the first version of this paper. We also thank the anonymous reviewers for their helpful suggestions.

Impact Statement

This paper presents work whose goal is to advance the field of Machine Learning and Quantum Computing. There are many potential societal consequences of our work, none which we feel must be specifically highlighted here.

References

- Aharonov, D. and Ta-Shma, A. Adiabatic quantum state generation and statistical zero knowledge. In *Proceedings of the 35th Annual ACM Symposium on Theory of Computing (STOC 2003)*, pp. 20–29, 2003.
- Ahn, S., Korattikara, A., and Welling, M. Bayesian posterior sampling via stochastic gradient fisher scoring. In *Proceedings of the 29th International Conference on Machine Learning (ICML 2012)*, pp. 1771–1778, 2012.
- Apers, S. and Sarlette, A. Quantum fast-forwarding: Markov chains and graph property testing, 2019.
- Arunachalam, S., Havlicek, V., Nannicini, G., Temme, K., and Wocjan, P. Simpler (classical) and faster (quantum) algorithms for Gibbs partition functions. In *2021 IEEE International Conference on Quantum Computing and Engineering (QCE)*, pp. 112–122. IEEE, 2021.
- Brooks, S., Gelman, A., Jones, G., and Meng, X.-L. *Handbook of Markov Chain Monte Carlo*. Chapman and Hall/CRC, May 2011. ISBN 9780429138508. doi: 10.1201/b10905. URL <http://dx.doi.org/10.1201/b10905>.
- Bubeck, S., Eldan, R., and Lehec, J. Sampling from a log-concave distribution with projected langevin monte carlo, 2015.
- Buser, P. A note on the isoperimetric constant. *Annales scientifiques de l’École Normale Supérieure*, Ser. 4, 15(2):213–230, 1982. doi: 10.24033/asens.1426. URL <http://www.numdam.org/articles/10.24033/asens.1426/>.
- Chakrabarti, S., Childs, A. M., Hung, S.-H., Li, T., Wang, C., and Wu, X. Quantum algorithm for estimating volumes of convex bodies. *ACM Transactions on Quantum Computing*, 4(3), 2023. arXiv:1908.03903.
- Chandler, D. Introduction to modern statistical. *Mechanics. Oxford University Press, Oxford, UK*, 5:449, 1987.
- Cheeger, J. A Lower Bound for the Smallest Eigenvalue of the Laplacian, pp. 195–200. Princeton University Press, Princeton, 1971. ISBN 9781400869312. doi: doi:10.1515/9781400869312-013. URL <https://doi.org/10.1515/9781400869312-013>.
- Cheng, X., Chatterji, N. S., Bartlett, P. L., and Jordan, M. I. Underdamped Langevin MCMC: A non-asymptotic analysis. In *Proceedings of the 31st Conference On Learning Theory (COLT 2018)*, volume 75 of *Proceedings of Machine Learning Research*, pp. 300–323. PMLR, 06–09 Jul 2018.
- Childs, A. M., Li, T., Liu, J.-P., Wang, C., and Zhang, R. Quantum algorithms for sampling log-concave distributions and estimating normalizing constants. In *Advances in Neural Information Processing Systems*, volume 35, pp. 23205–23217, 2022.
- Cornelissen, A. and Hamoudi, Y. A sublinear-time quantum algorithm for approximating partition functions. In *Proceedings of the 2023 Annual ACM-SIAM Symposium on Discrete Algorithms (SODA 2023)*, pp. 1245–1264. Society for Industrial and Applied Mathematics, 2023. doi: 10.1137/1.9781611977554.ch46. URL <https://doi.org/10.1137/1.9781611977554.ch46>.
- Dalalyan, A. S. Theoretical guarantees for approximate sampling from smooth and log-concave densities. *Journal of the Royal Statistical Society Series B: Statistical Methodology*, 79(3):651–676, 2017a.
- Dalalyan, A. S. Further and stronger analogy between sampling and optimization: Langevin monte carlo and gradient descent. In *Conference on Learning Theory*, pp. 678–689. PMLR, 2017b.
- Durmus, A., Majewski, S., and Miasojedow, B. Analysis of langevin monte carlo via convex optimization. *Journal of Machine Learning Research*, 20(73):1–46, 2019. URL <http://jmlr.org/papers/v20/18-173.html>.
- Fan, J., Yuan, B., and Chen, Y. Improved dimension dependence of a proximal algorithm for sampling. In *The Thirty Sixth Annual Conference on Learning Theory*, pp. 1473–1521. PMLR, 2023.

- Ge, R., Lee, H., and Lu, J. Estimating normalizing constants for log-concave distributions: Algorithms and lower bounds. In *Proceedings of the 52nd Annual ACM SIGACT Symposium on Theory of Computing (STOC 2020)*, pp. 579–586, 2020.
- Gilyén, A., Su, Y., Low, G. H., and Wiebe, N. Quantum singular value transformation and beyond: exponential improvements for quantum matrix arithmetics. In *Proceedings of the 51st Annual ACM SIGACT Symposium on Theory of Computing (STOC 2019)*. ACM, 2019. doi: 10.1145/3313276.3316366. URL <https://doi.org/10.1145%2F3313276.3316366>.
- Goel, S. Modified logarithmic sobolev inequalities for some models of random walk. *Stochastic Processes and their Applications*, 114(1):51–79, 2004. ISSN 0304-4149. doi: <https://doi.org/10.1016/j.spa.2004.06.001>. URL <https://www.sciencedirect.com/science/article/pii/S0304414904000912>.
- Grover, L. K. Fixed-point quantum search. *Physical Review Letters*, 95(15), 2005. doi: 10.1103/physrevlett.95.150501. URL <https://doi.org/10.1103%2Fphysrevlett.95.150501>.
- Harrow, A. W. and Wei, A. Y. Adaptive quantum simulated annealing for Bayesian inference and estimating partition functions. In *Proceedings of the 31st Annual ACM-SIAM Symposium on Discrete Algorithms (SODA 2020)*, pp. 193–212, 2020.
- Jerrum, M., Sinclair, A., and Vigoda, E. A polynomial-time approximation algorithm for the permanent of a matrix with nonnegative entries. *J. ACM*, 51:671–697, 2004. URL <https://api.semanticscholar.org/CorpusID:47361920>.
- Lee, H., Risteski, A., and Ge, R. Beyond log-concavity: Provable guarantees for sampling multi-modal distributions using simulated tempering langevin monte carlo. *Advances in neural information processing systems*, 31, 2018.
- Li, R., Zha, H., and Tao, M. Sqrt(d) dimension dependence of langevin monte carlo. In *The International Conference on Learning Representations*, 2022.
- Li, T. and Zhang, R. Quantum speedups of optimizing approximately convex functions with applications to logarithmic regret stochastic convex bandits. *Advances in Neural Information Processing Systems*, 35:3152–3164, 2022.
- Ma, Y.-A., Chatterji, N., Cheng, X., Flammarion, N., Bartlett, P., and Jordan, M. I. Is there an analog of nesterov acceleration for mcmc?, 2019a.
- Ma, Y.-A., Chen, Y., Jin, C., Flammarion, N., and Jordan, M. I. Sampling can be faster than optimization. *Proceedings of the National Academy of Sciences*, 116(42): 20881–20885, 2019b.
- Magniez, F., Nayak, A., Roland, J., and Santha, M. Search via quantum walk. *SIAM Journal on Computing*, 40(1): 142–164, 2011. doi: 10.1137/090745854. URL <https://doi.org/10.1137%2F090745854>.
- Montanaro, A. Quantum speedup of monte carlo methods. *Proceedings of the Royal Society A: Mathematical, Physical and Engineering Sciences*, 471(2181):20150301, sep 2015. doi: 10.1098/rspa.2015.0301. URL <https://doi.org/10.1098%2Frspa.2015.0301>.
- Murphy, K. P. *Probabilistic machine learning: an introduction*. MIT press, 2022.
- Neal, R. M. et al. Mcmc using hamiltonian dynamics. *Handbook of Markov chain Monte Carlo*, 2(11):2, 2011.
- Nesterov, Y. *Lectures on Convex Optimization*. Springer Optimization and Its Applications. Springer International Publishing, 2018. ISBN 9783319915784. URL <https://books.google.com/books?id=IPh6DwAAQBAJ>.
- Otto, F. and Villani, C. Generalization of an inequality by talagrand and links with the logarithmic sobolev inequality. *Journal of Functional Analysis*, 173(2):361–400, 2000. ISSN 0022-1236. doi: <https://doi.org/10.1006/jfan.1999.3557>. URL <https://www.sciencedirect.com/science/article/pii/S0022123699935577>.
- Pardo, L. *Statistical inference based on divergence measures*. CRC press, 2018.
- Raginsky, M., Rakhlin, A., and Telgarsky, M. Non-convex learning via stochastic gradient langevin dynamics: a nonasymptotic analysis. In *Conference on Learning Theory*, pp. 1674–1703. PMLR, 2017.
- Richter, P. C. Quantum speedup of classical mixing processes. *Physical Review A*, 76(4):042306, 2007.
- Shen, R. and Lee, Y. T. The randomized midpoint method for log-concave sampling. *Advances in Neural Information Processing Systems*, 32, 2019.
- Somma, R., Boixo, S., and Barnum, H. Quantum simulated annealing, 2007.
- Somma, R. D., Boixo, S., Barnum, H., and Knill, E. Quantum simulations of classical annealing processes. *Physical Review Letters*, 101(13), September 2008. ISSN 1079-7114. doi: 10.1103/physrevlett.101.130504. URL <http://dx.doi.org/10.1103/PhysRevLett.101.130504>.

- Stefankovic, D., Vempala, S., and Vigoda, E. Adaptive simulated annealing: A near-optimal connection between sampling and counting. *Journal of the ACM (JACM)*, 2009.
- Szegedy, M. Quantum speed-up of markov chain based algorithms. In *Proceedings of the 45th Annual IEEE Symposium on Foundations of Computer Science (FOCS 2004)*, pp. 32–41, 2004. doi: 10.1109/FOCS.2004.53.
- Vempala, S. Geometric random walks: a survey. In *Combinatorial and Computational Geometry*, pp. 577–616. Cambridge University Press, 2007.
- Vempala, S. and Wibisono, A. Rapid convergence of the unadjusted langevin algorithm: Isoperimetry suffices. *Advances in Neural Information Processing Systems*, 32, 2019.
- Wocjan, P. and Abeyesinghe, A. Speedup via quantum sampling. *Physical Review A*, 78(4):042336, 2008.
- Wocjan, P., Chiang, C.-F., Nagaj, D., and Abeyesinghe, A. Quantum algorithm for approximating partition functions. *Physical Review A*, 80(2):022340, 2009. arXiv:0811.0596.
- Xu, P., Chen, J., Zou, D., and Gu, Q. Global convergence of langevin dynamics based algorithms for nonconvex optimization. *Advances in Neural Information Processing Systems*, 31, 2018.
- Zhang, Y., Liang, P., and Charikar, M. A hitting time analysis of stochastic gradient Langevin dynamics. In *Conference on Learning Theory*, pp. 1980–2022. PMLR, 2017.
- Zou, D., Xu, P., and Gu, Q. Faster convergence of stochastic gradient Langevin dynamics for non-log-concave sampling. In *Uncertainty in Artificial Intelligence*, pp. 1152–1162. PMLR, 2021.

A. Appendix

B. Preliminaries

B.1. Classical MCMC

Monte Carlo Markov Chain (MCMC) methods are powerful computational techniques used for simulating and exploring complex probabilistic systems. In the context of sampling, MCMC involves constructing a Markov chain over the distribution's state space Ω , where each state represents a potential sample. By iteratively transitioning between states according to carefully designed transition probabilities, MCMC methods generate a sequence of samples that converge to the target distribution.

Let P be the transition matrix of a Markov chain over a finite state space Ω , and let π be the stationary distribution of this chain. A stationary distribution π is a probability distribution over the states that remain unchanged under the transition dynamics of the Markov chain. Mathematically, it satisfies the balance equation, $\pi^T P = \pi^T$. Hence π is left eigenvector of P with eigenvalue 1. The mixing time of a Markov chain can be defined as follows:

For any $0 < \epsilon < 1$, the mixing time $t_{\text{mix}}(\epsilon)$ is the smallest positive integer t such that for all initial distributions x in Ω :

$$\|P^t(x, \cdot) - \pi\|_{\text{TV}} \leq \epsilon. \quad (33)$$

Here, $P^t(x, \cdot)$ is the distribution of states after t steps starting from initial distribution x , and $\|\cdot\|_{\text{TV}}$ represents the total variation distance between two probability distributions. The mixing time $t_{\text{mix}}(\epsilon)$ characterizes the rate at which the chain approaches its stationary distribution within a specified tolerance ϵ . Bounding the mixing time of the Markov chain is often a primary obstacle when it comes to proving algorithm's total run time. Ergodicity and reversibility are crucial properties that significantly simplify the estimation of mixing time in Markov chains. A chain is said to be ergodic if it is irreducible (any state can be reached from any other state) and aperiodic (the chain does not return to the same state with periodic intervals). A Markov chain that satisfies the detailed balance condition is known as time-reversible and it is a fundamental property for establishing the mixing time in terms of spectral gap or conductance. Mathematically, it can be expressed as:

$$P_{xy}\pi_y = P_{yx}\pi_x. \quad (34)$$

This condition also guarantees that the Markov chain will converge to a stationary distribution. Furthermore, for a reversible Markov chain, it holds that

$$t_{\text{mix}}(\epsilon) \leq \frac{1}{\gamma} \cdot \log \left(\frac{1}{\pi^* \epsilon} \right). \quad (35)$$

where $\pi^* = \min_{x \in \Omega} \pi(x)$ and γ is the spectral gap defined as the difference between the first and the second-largest eigenvalue (in absolute value) of the transition matrix.

B.2. Langevin Diffusion

Langevin diffusion, often referred to as Langevin dynamics, is a fundamental stochastic differential equation that describes the dynamics of a particle undergoing random motion in a fluid or a complex environment. It is widely used in various scientific disciplines, including physics, chemistry, and biology, to model systems exhibiting Brownian motion or other forms of random behavior. It can be expressed as a continuous-time stochastic process X_t in the following form:

$$dX_t = -\nabla f(X_t) dt + \sqrt{2} dW_t, \quad (36)$$

where W_t is the standard Brownian motion.

Langevin diffusion provides a probabilistic approach to optimization by simulating the motion of particles under the influence of both deterministic gradient forces and random noise. This allows the optimization process to explore the parameter space more extensively, potentially escaping local optima and reaching a broader range of solutions. By simulating Langevin dynamics, machine learning practitioners can sample from the posterior distribution of the model parameters, enabling Bayesian inference and uncertainty estimation.

B.3. Unadjusted Langevin Algorithm (ULA)

Under certain conditions on f , Equation (36) accepts $e^{-f(x)}$ as its stationary density. Therefore, it is natural to discretize the Langevin diffusion using step size $\eta > 0$. The simplest discretization scheme is known as Euler-Maruyama method and it gives the following update rule:

$$\mathbf{x}_{k+1} = \mathbf{x}_k - \eta \nabla f(\mathbf{x}_k) + \sqrt{2\eta\beta^{-1}} \mathbf{z}_k, \quad (37)$$

where $\mathbf{z}_{\{0,1,\dots\}}$ are i.i.d Gaussian random vectors in $\mathbb{R}^d$.

Algorithm 1 Unadjusted Langevin Algorithm (ULA)

Input: $\mathbf{x}_0, (\eta > 0)$
Output: $\mathbf{x}_K$
for $k = 1, \dots, K$ **do**
 $\mathbf{x}_k = \mathbf{x}_{k-1} - \eta \nabla f(\mathbf{x}_{k-1}) + \sqrt{2\eta/\beta} \mathbf{z}_{k-1}$
end for
 Return $\mathbf{x}_K$

B.4. Metropolis Adjusted Langevin Algorithm (MALA)

Although ULA algorithm seems appealing due to its simplicity, it comes with a catch. Due to naive discretization of a continuous differential equation, the chain is asymptotically biased. That is, its stationary distribution is different than the stationary distribution of Langevin equation where the size of discrepancy depends on the step size, feature dimension and properties of f . Furthermore, the chain does not satisfy the detailed balance condition, which is the standard assumption in the mixing time analysis of Markov chains. A common practice is to apply Metropolis-Hasting correction at the end of each step to make the Markov chain reversible. This way, the chain converges to the desired target state and becomes time reversible. The update of the algorithm is modified such that if the following condition is satisfied, the iterates stay the same rather than applying Equation (37).

$$\frac{p(\mathbf{x}_k|\mathbf{x}_{k+1})\pi(\mathbf{x}_k)}{p(\mathbf{x}_{k+1}|\mathbf{x}_k)\pi(\mathbf{x}_{k+1})} < u, \quad (38)$$

where $u \sim \mathcal{U}[0, 1]$. This algorithm is called Metropolis adjusted Langevin algorithm (MALA).

Algorithm 2 Metropolis Adjusted Langevin Algorithm (MALA)

Input: $\mathbf{x}_0, (\eta > 0)$
Output: $\mathbf{x}_K$
for $k = 1, \dots, K$ **do**
 $\mathbf{x}_k = \mathbf{x}_{k-1} - \eta \nabla f(\mathbf{x}_{k-1}) + \sqrt{2\eta\beta^{-1}} \mathbf{z}_{k-1}$
 $\alpha = \frac{p(\mathbf{x}_{k-1}|\mathbf{x}_k)\pi(\mathbf{x}_{k-1})}{p(\mathbf{x}_k|\mathbf{x}_{k-1})\pi(\mathbf{x}_k)}$
 $u \sim \mathcal{U}[0, 1]$
 if $\alpha < u$ **then**
 $\mathbf{x}_k = \mathbf{x}_{k-1}$
 end if
end for
 Return $\mathbf{x}_K$

However, in many optimization problems, computing the gradient and applying Metropolis step can become highly costly in terms of computation time. For example, in large-scale machine learning the objection function consists of a sum with a large number of terms and MALA requires $\Omega(N)$ function evaluations for data size N . Motivated by the large-scale optimization problems, we focus on unadjusted version of the algorithm with stochastic gradients and we believe it is worthwhile to study the possible speed-up that could be achieved on a quantum computer.

C. Proofs for Annealing Schedule

We first restate the following useful lemmas from previous works as we use them repeatedly in our proofs. The first one lower bounds $f(x)$ by a quadratic function whereas the second one upper bounds the norm of the gradient by a linear function. We refer the readers to the original papers for their proofs and we don't repeat here for readability.

Lemma C.1 (Lemma A.1 in (Zou et al., 2021)). *Under Assumption 1.2, the objective function $f(x)$ satisfies,*

$$f(x) \geq \frac{m}{4} \|x\|^2 + f(x^*) - b/2, \quad (39)$$

where $f(x^*) = \min_x f(x)$.

Lemma C.2 (Lemma 3.1 in (Raginsky et al., 2017)). *Under Assumption 1.1, there exists a constant $G = \max_{k \in [N]} \|\nabla f_k(0)\|$ such that for any $x \in \mathbb{R}^d$ and $k \in [n]$, it holds that,*

$$\|\nabla f_k(x)\| \leq L\|x\| + G. \quad (40)$$

The next three technical lemmas are presented to make the proof of the annealing schedule concise. From more technical perspective, these lemmas generalizes the work done in (Ge et al., 2020) for non-logconcave distributions under Assumption 1.1 and Assumption 1.2.

Lemma C.3. *Suppose $\pi(x) \propto e^{-f(x)}$ is a Gibbs measure and f satisfies Assumptions 1.1 and 1.2. Then, we have*

$$\mathbb{E}_\pi [\exp(-s\|x\|^2)] \mathbb{E}_\pi [\exp(s\|x\|^2)] \leq O(\exp(dLs^2/(mc_s^2))), \quad (41)$$

where c_s^2 is the log-Sobolev constant of the distribution $\pi_s \propto \pi e^{s\|x\|^2}$.

Proof. Let $h(s) = \mathbb{E}_\pi [\exp(-s\|x\|^2)] \mathbb{E}_\pi [\exp(s\|x\|^2)]$, then

$$\frac{h'(s)}{h(s)} = \left(\frac{\mathbb{E}_\pi [\|x\|^2 \exp(s\|x\|^2)]}{\mathbb{E}_\pi [\exp(s\|x\|^2)]} - \frac{\mathbb{E}_\pi [\|x\|^2 \exp(-s\|x\|^2)]}{\mathbb{E}_\pi [\exp(-s\|x\|^2)]} \right) \quad (42)$$

$$= \int_{-s}^s v'(t) dt, \quad (43)$$

where $v(t)$ is defined as,

$$v(t) = \frac{\mathbb{E}_\pi [\|x\|^2 \exp(t\|x\|^2)]}{\mathbb{E}_\pi [\exp(t\|x\|^2)]}. \quad (44)$$

Computing $v'(t)$ gives,

$$v'(t) = \frac{\mathbb{E}_\pi [\|x\|^4 \exp(t\|x\|^2)] \mathbb{E}_\pi [\exp(t\|x\|^2)] - (\mathbb{E}_\pi [\|x\|^2 \exp(t\|x\|^2)]^2}{(\mathbb{E}_\pi [\exp(t\|x\|^2)])^2} \quad (45)$$

$$= \text{Var}_{\pi_t} \|x\|^2, \quad (46)$$

where π_t is a distribution defined as,

$$\pi_t(x) \propto \pi(x) \exp(t\|x\|^2). \quad (47)$$

Suppose π satisfies the log-Sobolev inequality with constant c_{LSI} , it also satisfies the Poincare inequality with the same constant (e.g (Goel, 2004)).

$$\text{Var}_{\pi_t} [\|x\|^2] \leq \frac{1}{c_t} \mathbb{E}_{\pi_t} [\|x\|^2] \leq O(Ld/(mc_t^2)), \quad (48)$$

where c_t is LSI constant of π_t and the second inequality is due to Lemma C.5. Therefore,

$$\frac{h'(s)}{h(s)} = \int_{-s}^s v'(t) dt = O(dLs/(mc_s^2)). \quad (49)$$

Hence,

$$\log(h(s)) - \log(h(0)) = \int_0^s \frac{h'(t)}{h(t)} dt = O(dLs^2/(mc_s^2)). \quad (50)$$

Since $h(0) = 1$, we conclude the proof. $\square$

Lemma C.4. Suppose $\pi(x) \propto e^{-f(x)}$ is a Gibbs measure and f satisfies the log-Sobolev inequality with constant c_{LSI} . Then under Assumptions 1.1 and 1.2,

$$\frac{\mathbb{E}_\pi [\exp(-(1+\alpha)\|x\|^2)] \mathbb{E}_\pi [\exp(-(1-\alpha)\|x\|^2)]}{(\mathbb{E}_\pi [\exp(-\|x\|^2)])^2} \leq O(\exp(dL\alpha^2/(c_{\text{LSI}}^2 m))) \quad (51)$$

for $0 \leq \alpha \leq 1/2$.

Proof. This follows from Lemma C.3, by setting $\tilde{\pi} \propto \pi \exp(-\|x\|^2)$. Then,

$$\frac{\mathbb{E}_\pi [\exp(-(1+\alpha)\|x\|^2)] \mathbb{E}_\pi [\exp(-(1-\alpha)\|x\|^2)]}{(\mathbb{E}_\pi [\exp(-\|x\|^2)])^2} = \mathbb{E}_{\tilde{\pi}} [\exp(-\alpha\|x\|^2)] \mathbb{E}_{\tilde{\pi}} [\exp(\alpha\|x\|^2)] \quad (52)$$

$$\leq O(\exp(dL\alpha^2/(mc_\alpha^2))) \quad (53)$$

$$\leq O(\exp(dL\alpha^2/(mc_{\text{LSI}}^2))) \quad (54)$$

with c_α is LSI constant of $\pi_\alpha \propto \pi \exp(-(1-\alpha)\|x\|^2)$. The last step follows from the fact that $c_\alpha \geq c_{\text{LSI}}$ for $\alpha \leq 1/2$. $\square$

Lemma C.5. Suppose $\pi(x) \propto e^{-f(x)}$ is a Gibbs measure and f satisfies Assumptions 1.1 and 1.2. Then

$$\mathbb{E}_{\pi_s(x)} [e^{s\|x\|^2} \|x\|^2] \leq O(Ld/(mc_s)), \quad (55)$$

where π_s is a probability distribution proportional to $\pi(x)e^{s\|x\|^2}$ for a constant $s \leq \frac{m}{8}$ and c_s is the log-Sobolev constant of π_s .

Proof. Our proof follows the idea presented in proof of Lemma 6 in (Ma et al., 2019b) without the assumption of local non-convexity. We choose an auxiliary random variable x' following the law of $p \propto e^{-(L-s)\|x\|^2}$ and couples optimally with $x_s \sim \pi_s : (x_s, x') \sim \gamma \in \Gamma_{\text{opt}}(\pi_s, p)$.

$$\mathbb{E}_{\pi_s} \|x\|^2 = \mathbb{E}_{(x_s, x' \sim \gamma)} \|x' - x' + x_s\|^2 \quad (56)$$

$$\leq 2\mathbb{E}_p \|x'\|^2 + 2\mathbb{E}_{(x_s, x' \sim \gamma)} \|x' - x_s\|^2 \quad (57)$$

$$= \frac{2d}{L-s} + 2W_2^2(p, \pi_s) \quad (58)$$

$$\leq \frac{2d}{L-s} + \frac{2}{c_s} \text{KL}(p, \pi_s), \quad (59)$$

where c_{π_s} is LSI constant of π_s . The first inequality follows from Young's inequality and second inequality is due to generalized Talagrand inequality (Otto & Villani, 2000). KL divergence can be bounded,

$$\text{KL}(p, \pi_s) = \int_x dx \log \left(\frac{p(x)}{\pi_s(x)} \right) p(x) \quad (60)$$

$$\leq \sup_x \log \left(\frac{p(x)}{\pi_s(x)} \right) \int_x dx p(x) \quad (61)$$

$$= \sup_x \log \left(\frac{p(x)}{\pi_s(x)} \right). \quad (62)$$

We can further bound $\frac{p(x)}{\pi_s(x)}$ for any $x \in \Omega$,

$$\frac{p(x)}{\pi_s(x)} = \frac{e^{-(L-s)\|x\|^2}}{\int_x dx e^{-(L-s)\|x\|^2}} \frac{\int dx e^{-f(x)} e^{s\|x\|^2}}{e^{-f(x)} e^{s\|x\|^2}} \quad (63)$$

$$= \frac{\int dx e^{s\|x\|^2 - f(x)}}{\int dx e^{-(L-s)\|x\|^2}} e^{-L\|x\|^2 + f(x)} \quad (64)$$

$$\leq \frac{\int dx e^{s\|x\|^2 - m\|x\|^2/4 + b/2 - f(x^*)}}{\int dx e^{-(L-s)\|x\|^2}} e^{-L\|x\|^2 + f(x)} \quad (65)$$

$$\leq \frac{\int dx e^{s\|x\|^2 - m\|x\|^2/4 + b/2 - f(x^*)}}{\int dx e^{-(L-s)\|x\|^2}} e^{L\|x^*\|^2 + f(x^*)} \quad (66)$$

$$= e^{b/2 + L\|x^*\|^2} \frac{(L-s)^{d/2}}{(m/4 - s)^{d/2}}, \quad (67)$$

where the first inequality is due to Assumption 1.2 and Lemma C.1. Second inequality follows from Equation (70). Hence, KL divergence is bounded by,

$$\text{KL}(p, \pi_s) \leq \sup_x \left(\frac{p(x)}{\pi_s(x)} \right) \leq b/2 + L\|x^*\|^2 + \frac{d}{2} \log \left(\frac{L-s}{m/2 - 2s} \right). \quad (68)$$

This implies that,

$$\mathbb{E}_{\pi_s} \|x\|^2 \leq \frac{2d}{L-s} + \frac{2}{c_s} (b/2 + L\|x^*\|^2 + \frac{d}{2} \log \left(\frac{L-s}{m/2 - 2s} \right)) = O(Ld/(mc_s)) \quad (69)$$

for $s \leq m/8$. $\square$

Finally we are ready to prove our result for quantum annealing procedure. The key idea in this proofs is to show that two consequent Gibbs distributions in our annealing scheme are close to each other so that the Markov chains become slowly changing. Furthermore, we show that the final distribution is close to desired Gibbs distribution.

Lemma 3.1 (Quantum Annealing). *Under Assumptions 1.1 and 1.2, there exists a series of quantum states $|\mu_0\rangle, |\mu_1\rangle, \dots, |\mu_M\rangle$ satisfying the following properties:*

1. *There exists an efficient quantum algorithm to prepare initial state $|\mu_0\rangle$ without using any function queries.*

2. *For all $i \in \{0, \dots, M-1\}$, $|\mu_i\rangle$ and $|\mu_{i+1}\rangle$ has at least constant overlap, i.e.,*

$$|\langle \mu_i | \mu_{i+1} \rangle| \geq \Omega(1). \quad (12)$$

3. *The final state $|\mu_M\rangle$ has at least constant overlap with the target Gibbs state $|\pi\rangle$,*

$$|\langle \mu_M | \pi \rangle| \geq \Omega(1). \quad (13)$$

4. *The number of quantum states $M \leq \tilde{O}(c_{\text{LSI}}^{-1} d^{1/2})$.*

Proof. Our construction and analysis are similar to the annealing scheme used in (Childs et al., 2022), however our proof does not require any convexity assumption for $f(x)$. The construction is as follows:

1. $|\mu_0\rangle = \sum_x \sqrt{p_0(x)} |x\rangle$, where $p_0(x) = \frac{\exp\left(-\frac{\|x\|^2}{2\sigma_1^2}\right)}{Z_0}$.
2. For all $i \in [1, M-1]$, $|\mu_i\rangle = \sum_{x \in \Omega} \sqrt{p_i(x)} |x\rangle$, where $p_i(x) = \frac{\exp\left(-f(x) - \frac{\|x\|^2}{2\sigma_i^2}\right)}{Z_i}$ such that $\sigma_{i+1}^2 = \sigma_i^2(1 + \alpha)$ with $\alpha = \tilde{O}(d^{-1/2} c_{\text{LSI}})$.

Here, $Z_0 = \int dx \exp\left(-\frac{\|x\|^2}{2\sigma_1^2}\right)$ and $Z_i = \int dx \exp\left(-f(x) - \frac{\|x\|^2}{2\sigma_i^2}\right)$. The first property in the lemma statement holds, since p_0 corresponds to a Gaussian distribution and the coherent quantum state corresponding to Gaussian distributions can be efficiently prepared by using Box-Muller technique without using any evaluation of f or ∇f . Next, we prove the second property. We first start with $i = 0$ as the base case: $\langle \mu_0 | \mu_1 \rangle \geq \Omega(1)$. To prove this, Let $f(x^*) = \min_{x \in \Omega} f(x)$. We fix $\beta = 1$ without loss of generality. Then, we can write,

$$f(x) \leq f(x^*) + \langle \nabla f(x^*), x - x^* \rangle + \frac{L}{2} \|x - x^*\|^2 \leq f(x^*) + L\|x^*\|^2 + L\|x\|^2, \quad (70)$$

where the first inequality is well known due to Assumption 1.1 (see (Nesterov, 2018)) and second inequality is due to Young's inequality. Using this upper bound on $f(x)$, we have,

$$|\langle \mu_0 | \mu_1 \rangle| = \frac{\int dx \exp\left(-\frac{1}{2}f(x) - \frac{\|x\|^2}{2\sigma_1^2}\right)}{(2\pi\sigma_1^2)^{d/4} \sqrt{Z_1}} \quad (71)$$

$$\geq \frac{\int dx \exp\left(-\frac{1}{2}f(x^*) - \frac{1}{2}L\|x\|^2 - \frac{1}{2}L\|x^*\|^2 - \frac{\|x\|^2}{2\sigma_1^2}\right)}{(2\pi\sigma_1^2)^{d/4} \sqrt{\int dx \exp\left(-f(x^*) - \frac{\|x\|^2}{4\sigma_1^2}\right)}} \quad (72)$$

$$= \frac{\exp\left(-\frac{L}{2}\|x^*\|^2\right) \pi^{d/2} (L/2 + 1/(2\sigma_1^2))^{-d/2}}{(2\pi\sigma_1^2)^{d/4} (2\pi\sigma_1^2)^{d/4}} \quad (73)$$

$$= \exp\left(-\frac{L}{2}\|x^*\|^2\right) (L\sigma_1^2 + 1)^{-d/2} \quad (74)$$

$$\geq \exp\left(-\frac{L}{2}\|x^*\|^2 - \frac{dL\sigma_1^2}{2}\right). \quad (75)$$

Choosing $\sigma_1^2 = \frac{\epsilon}{2dL}$ yields $|\langle \mu_0 | \mu_1 \rangle| \geq \Omega(1)$. Next, we consider $1 \leq i \leq M-1$. Letting $\sigma^2 = \sigma_{i+1}^2$, we have

$$|\langle \mu_i | \mu_{i+1} \rangle| = \int dx \frac{\exp(-f_i(x)/2)}{\sqrt{Z_i}} \frac{\exp(-f_{i+1}(x)/2)}{\sqrt{Z_{i+1}}} \quad (76)$$

$$= \int dx \frac{\exp\left(-f(x) - \frac{\|x\|^2}{4\sigma_i^2} - \frac{\|x\|^2}{4\sigma_{i+1}^2}\right)}{\sqrt{Z_i Z_{i+1}}} \quad (77)$$

$$= \frac{\mathbb{E}_\pi \left[\exp\left(-\frac{1+\alpha/2}{2\sigma^2} \|x\|^2\right) \right]}{\mathbb{E}_\pi \left[\exp\left(-\frac{1+\alpha}{2\sigma^2} \|x\|^2\right) \right]^{1/2} \mathbb{E}_\pi \left[\exp\left(-\frac{1-\alpha}{2\sigma^2} \|x\|^2\right) \right]^{1/2}}, \quad (78)$$

where the last step follows from the fact that the numerator can be written as,

$$\int dx \exp\left(-f(x) - \frac{\|x\|^2}{4\sigma_i^2} - \frac{\|x\|^2}{4\sigma_{i+1}^2}\right) = \frac{Z \int dx \exp\left(-f(x) - \frac{1+\alpha/2}{2\sigma^2} \|x\|^2\right)}{Z} = Z \mathbb{E}_\pi \left[\exp\left(-\frac{1+\alpha/2}{2\sigma^2} \|x\|^2\right) \right], \quad (79)$$

and similarly, Z_i and Z_{i+1} can be simplified as,

$$Z_i = \int dx e^{-f(x) - \frac{1}{2\sigma_i^2} \|x\|^2} = \frac{Z \int dx \exp\left(-f(x) - \frac{(1+\alpha)}{2\sigma^2} \|x\|^2\right)}{Z} = Z \mathbb{E}_\pi \left[\exp\left(-\frac{(1+\alpha)}{2\sigma^2} \|x\|^2\right) \right] \quad (80)$$

$$Z_{i+1} = \int dx e^{-f(x) - \frac{1}{2\sigma_{i+1}^2} \|x\|^2} = \frac{Z \int dx \exp\left(-f(x) - \frac{1-\alpha}{2\sigma^2} \|x\|^2\right)}{Z} = Z \mathbb{E}_\pi \left[\exp\left(-\frac{1-\alpha}{2\sigma^2} \|x\|^2\right) \right]. \quad (81)$$

Defining $\alpha' = \frac{\alpha}{\alpha+2}$ and $\sigma'^2 = \frac{\sigma^2}{1+\alpha/2}$, we have

$$|\langle \mu_i | \mu_{i+1} \rangle| = \frac{\mathbb{E}_\pi \left[\exp\left(-\frac{1}{2\sigma'^2} \|x\|^2\right) \right]}{\mathbb{E}_\pi \left[\exp\left(-\frac{1+\alpha'}{2\sigma'^2} \|x\|^2\right) \right]^{1/2} \mathbb{E}_\pi \left[\exp\left(-\frac{1-\alpha'}{2\sigma'^2} \|x\|^2\right) \right]^{1/2}} \quad (82)$$

$$\geq \Omega(\exp(-2dL\alpha'^2/(mc_{\text{LSI}}^2))), \quad (83)$$

where the last inequality is due to *Lemma C.4*. Setting $\alpha^2 = \tilde{O}(c_{\text{LSI}}^2 m/(dL))$, we have $|\langle \mu_i | \mu_{i+1} \rangle| \geq \Omega(1)$. Having established the second property, we move on to the third property.

$$|\langle \mu_M | \pi \rangle| = \int dx \frac{\exp\left(-f(x) - \frac{\|x\|^2}{4\sigma_M^2}\right)}{\sqrt{Z_M} \sqrt{Z}} \quad (84)$$

$$= \mathbb{E}_{\rho'} \left[\exp\left(-\frac{1}{4\sigma_M^2} \|x\|^2\right) \right]^{-1/2} \mathbb{E}_{\rho'} \left[\exp\left(\frac{1}{4\sigma_M^2} \|x\|^2\right) \right]^{-1/2} \quad (85)$$

$$\geq 1 - \Omega(dL/(m\sigma_M^4 c_{\text{LSI}}^2)), \quad (86)$$

where $\rho' \propto \pi(x) \exp\left(-\frac{\|x\|^2}{4\sigma_M^2}\right)$. The last step is due to *Lemma C.3*. Setting $\sigma_M^2 = \sqrt{dL/(mc_{\text{LSI}}^2)}$ satisfies $|\langle \mu_M | \pi \rangle| \geq \Omega(1)$. The final property follows from the fact that $\alpha = \tilde{O}(\sqrt{c_{\text{LSI}}^2 m/(dL)})$, since,

$$\sigma_M = \sigma_0(1 + \alpha)^M, \quad (87)$$

and solving this for M yields $M = \tilde{O}(\sqrt{dL/(mc_{\text{LSI}}^2)})$. $\square$

D. Proofs for Quantum Sampling Algorithms

For technical reasons, we set the domain $\Omega = \mathbb{R}^d \cap B(0, R)$ where R is sufficiently large enough to show that the truncated distribution π^* in Ω is ϵ close to the original Gibbs distribution. More specifically, we work on sufficiently large but bounded domain to show that the norm of the gradients are bounded and derive our results in terms of R . The truncation is done by only considering the sum of projectors up to $\|x\| \leq R$ in the implementation of the quantum walk. Then we use the following lemma to characterize R ,

Lemma D.1 (Lemma 6 in (Zou et al., 2021)). *For any $\epsilon \in (0, 1)$ set $R = \bar{R}(\epsilon/12)$ and let π^* be the truncated distribution in Ω . Then the total variation distance between π^* and π is upper bounded by $\|\pi^* - \pi\| \leq \epsilon/4$, where*

$$\bar{R}(z) = \left[\max \left\{ \frac{625d \log(4/z)}{m\beta}, \frac{4d \log(4L/m)}{m\beta}, \frac{4d + 8\sqrt{d \log(1/z)} + 8 \log(1/z)}{m\beta} \right\} \right]^{1/2}. \quad (88)$$

D.1. Proofs for Quantum MALA

The lemma below characterizes the conductance parameter of the classical MALA algorithm constructed with stochastic gradients under given assumptions. Though similar results are given for full gradient case in (Ma et al., 2019b), we use the stochastic version and remove B dependent condition on the step size when we apply this lemma in full gradient case by setting $B \gg d$.

Lemma D.2 (Lemma 6.5 in (Zou et al., 2021)). *Under Assumptions 1.1 and 1.2, if the step size meets the condition $\eta \leq \min \{35(Ld + (LR + G)^2 \beta d/B)^{-1}, [25\beta(LR + G)^2]^{-1}\}$, then there exists absolute constant c_0 such that, the conductance parameter ϕ for Metropolis adjusted Stochastic Langevin Algorithm satisfies,*

$$\phi \geq c_0 \rho \sqrt{\eta/\beta}, \quad (89)$$

where ρ is the Cheeger constant of the truncated distribution π^* .

The following lemma is useful to characterize the phase gap of quantum walk operator for a reversible Markov chain in terms of its conductance parameter and it is the source of the quantum speed up for mixing time for reversible chains.

Lemma D.3. *Let Q be a reversible Markov chain with conductance parameter $\phi(Q)$ and let eigenvalues for the transition density of Q be $\lambda_0 = 1 > |\lambda_1| \geq |\lambda_2| \geq \dots \geq |\lambda_m|$. Let W be a unitary quantum walk operator constructed with the transition density of Q . Then the phase gap $\Delta(W) := 2 \arccos |\lambda_1|$ is lower bounded by,*

$$\Delta(W) \geq \sqrt{2} \phi(Q). \quad (90)$$

Proof. Let $\gamma(Q) = 1 - \lambda_1$ denote the spectral gap of Q . Using Cheeger's inequality (Cheeger, 1971), $\gamma(Q)$ can be bounded in terms of the conductance parameter,

$$\sqrt{2\gamma(Q)} \leq \phi(Q). \quad (91)$$

Let $\theta = \arccos |\lambda_1|$. Then we can write,

$$\Delta(W) \geq |1 - e^{2i\theta}| = 2\sqrt{1 - \lambda_1^2} \geq 2\sqrt{\gamma(Q)}. \quad (92)$$

By combining Equation (91) and Equation (92), we obtain $\Delta(Q) \geq \sqrt{2}\phi(Q)$. $\square$

Having established the phase gap of the quantum walk operator associated with quantum MALA algorithm, we are now ready to prove the following theorem.

Theorem 4.1 (Quantum MALA). *Let $\pi \propto e^{-\beta f(x)}$ denote a probability distribution with inverse temperature $\beta > 0$ such that $f(x)$ satisfies Assumptions 1.1 and 1.2. Then, there exists a quantum algorithm that outputs a random variable distributed according to μ such that,*

$$\|\mu - \pi\|_{\text{TV}} \leq \epsilon, \quad (26)$$

where $\|\cdot\|_{\text{TV}}$ is the total variation distance, using $\tilde{O}(\beta d \rho^{-1} c_{\text{LSI}}^{-1})$ queries to $\mathcal{O}_{\nabla f}$ and $\mathcal{O}_f$.

Proof. Let $|\mu_0\rangle, |\mu_1\rangle, \dots, |\mu_{M-1}\rangle$ be the series of quantum states described in Lemma 3.1. We start with the preparation of the initial Gaussian state $|\mu_0\rangle$ which can be done efficiently by applying the Box-Muller transformation to the uniform distribution state (see Appendix A.3 in Chakrabarti et al. (2023) for more details). Then, for each $i \in [0, M-2]$, we drive each state $|\mu_i\rangle$ to $|\mu_{i+1}\rangle$ using $\pi/3$ -fixed-point amplitude amplification algorithm (Grover, 2005). The amplitude amplification uses the following reflection operators,

$$V_i = e^{i\pi/3} |\mu_i\rangle \langle \mu_i| + (I - |\mu_i\rangle \langle \mu_i|), \quad (93)$$

$$V_{i+1} = e^{i\pi/3} |\mu_{i+1}\rangle \langle \mu_{i+1}| + (I - |\mu_{i+1}\rangle \langle \mu_{i+1}|). \quad (94)$$

Each state $|\mu_i\rangle$ is the unique eigenvector of quantum MALA operator U_i^* for $f_i(x) = f(x) + \frac{\|x\|^2}{2\sigma_i^2}$ since the classical MALA is time reversible and its stationary distribution is μ_i . Therefore, the operator $|\mu_i\rangle \langle \mu_i|$ is a projector operator to the eigenstate of U_i^* with eigenphase 0. Then, by Corollary 4.1 in (Chakrabarti et al., 2023), the operator V_i can be implemented with ϵ accuracy using $\tilde{O}(1/\Delta(U_i^*))$ calls to controlled- U^* operators where $\Delta(\cdot)$ is the phase gap. By Lemma D.2, Lemma D.1, and Lemma D.3, $\Delta \geq \rho\sqrt{2\eta/\beta}$ for step size smaller than $O(\min\{d^{-1}, \beta^{-1}\})$. Then, using $\tilde{O}(\rho^{-1}\eta^{-1/2}\beta)$ calls to controlled- U_i^* operators, we can implement a quantum reflection $\tilde{V}_i$ such that,

$$\|\tilde{V}_i - V_i\| \leq \epsilon, \quad (95)$$

for each i . Then, given $|\mu_i\rangle$, we can drive $|\mu_i\rangle$ to $|\tilde{\mu}_{i+1}\rangle$ using constant number of $\tilde{V}_i$ operators because,

$$|\langle \mu_i | \tilde{\mu}_{i+1} \rangle| \geq \Omega(1), \quad (96)$$

such that $\|\tilde{\mu}_{i+1}\rangle - |\mu_{i+1}\rangle\| \leq \epsilon$. Then we apply the same steps M times to drive μ_0 to π with at most error ϵ . Note that the error in each step does not accumulate linearly. This is because we can drive $\tilde{\mu}_i$ to $\tilde{\mu}_{i+1}$ with logarithmic cost in applying reflection operators. Since $M = \tilde{O}(c_{\text{LSI}}^{-1}\sqrt{d})$, the total complexity of the annealing procedure is $\tilde{O}(d^{1/2}c_{\text{LSI}}^{-1}\rho^{-1}\eta^{-1/2}\beta) = \tilde{O}(c_{\text{LSI}}^{-1}\rho^{-1}\beta d)$. Each U^* operator can be implemented using constant number of calls to full gradient and evaluation oracles, the algorithm uses $\tilde{O}(c_{\text{LSI}}^{-1}\rho^{-1}\beta d)$ full gradient and function evaluations. $\square$

D.2. Proofs for Quantum ULA

Since unadjusted Langevin algorithm is not a reversible chain, we cannot follow the same procedure since there is no direct relation between the conductance and phase gap. The following lemma, proved in Appendix D.2, quantifies the error between two quantum walk operators corresponding to different Markov chains in spectral norm.

Lemma D.4. Let U and $\tilde{U}$ be two quantum walk operators associated with two classical Markov chains with transition densities $P(x \rightarrow y) = p_{xy}$ and $\tilde{P}(x \rightarrow y) = \tilde{p}_{xy}$, respectively. Then,

$$\|U - \tilde{U}\| \leq 4\sqrt{2} \max_x \|P(x, \cdot) - \tilde{P}(x, \cdot)\|_{\text{H}}, \quad (97)$$

where $\|P(x, \cdot) - \tilde{P}(x, \cdot)\|_{\text{H}}$ denotes the Hellinger distance between the probability densities $P(x, \cdot)$ and $\tilde{P}(x, \cdot)$ for any $x \in \Omega$.

Proof. We first define the following quantum states,

$$|\psi_x\rangle = \sum_y \sqrt{p_{xy}} |x\rangle |y\rangle, \quad (98)$$

$$|\tilde{\psi}_x\rangle = \sum_y \sqrt{\tilde{p}_{xy}} |x\rangle |y\rangle. \quad (99)$$

Then, using the definition of quantum walk operators, the spectral norm of difference of operators can be bounded as,

$$\|U - \tilde{U}\| = \|S(2 \sum_{x \in \Omega} |\psi_x\rangle \langle \psi_x| - I) - S(2 \sum_{x \in \Omega} |\tilde{\psi}_x\rangle \langle \tilde{\psi}_x| - I)\| \quad (100)$$

$$\leq 2 \left\| \sum_{x \in \Omega} |\psi_x\rangle \langle \psi_x| - \sum_{x \in \Omega} |\tilde{\psi}_x\rangle \langle \tilde{\psi}_x| \right\| \quad (101)$$

$$\leq 2 \left\| \sum_{x \in \Omega} (|\psi_x\rangle - |\tilde{\psi}_x\rangle) \langle \psi_x| + \sum_{x \in \Omega} |\tilde{\psi}_x\rangle (\langle \psi_x| - \langle \tilde{\psi}_x|) \right\| \quad (102)$$

$$\leq 2 \left\| \sum_{x \in \Omega} (|\psi_x\rangle - |\tilde{\psi}_x\rangle) \langle \psi_x| \right\| + 2 \left\| \sum_{x \in \Omega} (|\psi_x\rangle - |\tilde{\psi}_x\rangle) \langle \tilde{\psi}_x| \right\|, \quad (103)$$

where the first inequality is due to unitarity of S and the third inequality is due to triangular inequality. Let $|\phi\rangle$ and $|\phi'\rangle$ are the states defined as the maximizers,

$$\left\| \sum_{x \in \Omega} (|\psi_x\rangle - |\tilde{\psi}_x\rangle) \langle \psi_x| \right\| = \max_{|\phi\rangle} \left\| \sum_{x \in \Omega} (|\psi_x\rangle - |\tilde{\psi}_x\rangle) \langle \psi_x | \phi \rangle \right\|, \quad (104)$$

and

$$\left\| \sum_{x \in \Omega} (|\psi_x\rangle - |\tilde{\psi}_x\rangle) \langle \tilde{\psi}_x| \right\| = \max_{|\phi'\rangle} \left\| \sum_{x \in \Omega} (|\psi_x\rangle - |\tilde{\psi}_x\rangle) \langle \tilde{\psi}_x | \phi' \rangle \right\|. \quad (105)$$

Notice that, for any $x \in \Omega$, we have $\langle \tilde{\psi}_x | \tilde{\psi}_y \rangle = \delta_{xy}$ and $\langle \psi_x | \psi_y \rangle = \delta_{xy}$. Therefore, we can write $|\phi\rangle = \sum_{x \in \Omega} c_x |\psi_x\rangle + |\xi\rangle$ and $|\phi'\rangle = \sum_{x \in \Omega} \tilde{c}_x |\tilde{\psi}_x\rangle + |\tilde{\xi}\rangle$ where $\langle \tilde{\xi} | \tilde{\psi}_x \rangle = \langle \xi | \psi_x \rangle = 0$ for all $x \in \Omega$. Hence,

$$\|U - \tilde{U}\| \leq 2 \left\| \sum_x c_x (|\tilde{\psi}_x\rangle - |\psi_x\rangle) \right\| + 2 \left\| \sum_x \tilde{c}_x (|\tilde{\psi}_x\rangle - |\psi_x\rangle) \right\| \quad (106)$$

$$\leq 4 \max_x \| |\tilde{\psi}_x\rangle - |\psi_x\rangle \|. \quad (107)$$

Finally, we can write,

$$\| |\tilde{\psi}_x\rangle - |\psi_x\rangle \| = \left\| \sum_y (\sqrt{p_{xy}} - \sqrt{\tilde{p}_{xy}}) |x\rangle |y\rangle \right\| \quad (108)$$

$$= \left(\sum_y (\sqrt{p_{xy}} - \sqrt{\tilde{p}_{xy}})^2 \right)^{1/2} \quad (109)$$

$$\leq \sqrt{2} \|P(x, \cdot) - \tilde{P}(x, \cdot)\|_{\text{H}}, \quad (110)$$

where the last step follows from the definition of Hellinger distance. $\square$

To be able to apply Lemma D.4, we bound the Hellinger distance between the probability distributions of MALA and ULA algorithm through the next lemma, which is proved in Appendix D.2.

Lemma D.5. *Let $P(x \rightarrow y) = p_{xy}$ and $P^*(x \rightarrow y) = p_{xy}^*$ be the transition densities for Unadjusted Langevin Algorithm (ULA) and Metropolis Adjusted Langevin Algorithm (MALA) respectively. Then under Assumptions 1.1 and 1.2 and for step size $\eta \leq d(\beta(LR + G)^2)^{-1}$,*

$$\max_x \|P(x, \cdot) - \tilde{P}(x, \cdot)\|_H \leq 4\eta dL, \quad (111)$$

where G is a positive constant that satisfies $\|\nabla f(0)\| \leq G$.

Proof. For the sake of the proof, we use the lazy version of the Markov chains as it does not change the stationary density.

Let $q_{xy} = \frac{1}{(4\pi\eta/\beta)^{d/2}} \exp\left(-\frac{\|y-x+\eta\nabla f(x)\|^2}{2\eta/\beta}\right)$, then we can write

$$p_{xy} = \frac{1}{2}\delta_{xy} + \frac{1}{2}q_{xy}, \quad (112)$$

and

$$p_{xy}^* = \begin{cases} \alpha_x(y)p_{xy}, & \text{if } x \neq y \\ p_{xy} + \sum_{z \in \Omega} p_{xz}(1 - \alpha_x(z)) & \text{if } x = y \end{cases}, \quad (113)$$

where δ_{xy} is Kronecker delta function and $\alpha_x(y)$ is the acceptance probability given by

$$\alpha_x(y) = \min \left(1, \frac{\exp\left(-\beta f(y) - \frac{\|x-y+\eta\nabla f(y)\|^2}{4\eta/\beta}\right)}{\exp\left(-\beta f(x) - \frac{\|y-x+\eta\nabla f(x)\|^2}{4\eta/\beta}\right)} \right). \quad (114)$$

By this definition, $\alpha_x(y) \leq 1$. Suppose $\alpha_x(y) \geq 1 - e(x, y)$. Then for $x \neq y$,

$$(\sqrt{p_{xy}^*} - \sqrt{p_{xy}})^2 = p_{xy}(1 - \sqrt{\alpha_{xy}})^2 \quad (115)$$

$$\leq p_{xy}(1 - \sqrt{1 - e(x, y)})^2 \quad (116)$$

$$\leq p_{xy}e(x, y)^2, \quad (117)$$

where the second inequality is due to the fact that for $0 \leq x \leq 1$,

$$1 - \sqrt{1 - x} = \frac{(1 - \sqrt{1 - x})(1 + \sqrt{1 - x})}{(1 + \sqrt{1 - x})} = \frac{1 - (1 - x)}{1 + \sqrt{1 - x}} = \frac{x}{1 + \sqrt{1 - x}} \leq x. \quad (118)$$

For $x = y$,

$$(\sqrt{p_{xy}^*} - \sqrt{p_{xy}})^2 = p_{xy} \left(\sqrt{1 + \frac{1 - \mathbb{E}_{p_{xy}}(\alpha_x(y))}{p_{xy}}} - 1 \right)^2 \quad (119)$$

$$\leq p_{xy} \left(1 + \frac{1 - \mathbb{E}_{p_{xy}}(\alpha_x(y))}{2p_{xy}} - 1 \right)^2 \quad (120)$$

$$\leq \frac{(1 - \mathbb{E}_{p_{xy}}(\alpha_x(y)))^2}{4p_{xy}} \quad (121)$$

$$\leq \frac{\mathbb{E}_{p_{xy}}(e(x, y))^2}{2}, \quad (122)$$

where the second inequality follows from $\sqrt{1+x} \leq 1 + \frac{x}{2}$ for $x \geq 0$ and the third inequality holds since $p_{xy} \geq \frac{1}{2}$ for $x = y$ because of laziness of the Markov chains. Therefore,

$$\int_{y \in \Omega} (\sqrt{p_{xy}^*} - \sqrt{p_{xy}})^2 dy = \int_{y \in \Omega} \delta_{xy} (\sqrt{p_{xy}^*} - \sqrt{p_{xy}})^2 dy + \int_{y \in \Omega} (1 - \delta_{xy}) (\sqrt{p_{xy}^*} - \sqrt{p_{xy}})^2 dy \quad (123)$$

$$\leq \mathbb{E}_{p_{xy}}(e(x, y)^2) + \frac{\mathbb{E}_{p_{xy}}(e(x, y))^2}{2} \quad (124)$$

$$\leq \frac{\mathbb{E}_{q_{xy}}(e(x, y)^2)}{2} + \frac{\mathbb{E}_{q_{xy}}(e(x, y))^2}{8}, \quad (125)$$

where the extra factors of $1/2$ and $1/4$ in the second inequality comes from the laziness of the chain. Now, we need to bound $e(x, y)$. Starting from

$$\alpha_x(y) \geq \frac{\exp\left(-\beta f(y) - \frac{\|x-y+\eta\nabla f(y)\|^2}{4\eta/\beta}\right)}{\exp\left(-\beta f(x) - \frac{\|y-x+\eta\nabla f(x)\|^2}{4\eta/\beta}\right)} \quad (126)$$

$$= \exp\left(-\beta(f(y) - f(x)) - \frac{2\eta\langle y-x, \nabla f(y) + \nabla f(x) \rangle + \eta^2\|\nabla f(y)\|^2 - \eta^2\|\nabla f(x)\|^2}{4\eta}\right) \quad (127)$$

$$\geq \exp\left(-\frac{\beta L\|x-y\|^2}{2} - \frac{\beta\eta^2\|\nabla f(y)\|^2 - \eta^2\|\nabla f(x)\|^2}{4\eta}\right) \quad (128)$$

$$\geq \exp\left(-\frac{\beta L\|x-y\|^2}{2} - \frac{\beta\eta L(LR+G)\|x-y\|}{2}\right) \quad (129)$$

$$\geq 1 - \frac{\beta L\|x-y\|^2}{2} - \frac{\beta\eta L(LR+G)\|x-y\|}{2}. \quad (130)$$

The second inequality holds because of the smoothness of $f(x)$ since,

$$f(x) \leq f(y) + \langle y-x, \nabla f(x) \rangle + \frac{L\|x-y\|^2}{2}, \quad (131)$$

$$f(y) \leq f(x) + \langle x-y, \nabla f(y) \rangle + \frac{L\|x-y\|^2}{2}, \quad (132)$$

which implies the following inequality

$$|f(y) - f(x) - \frac{1}{2}\langle y-x, \nabla f(x) + \nabla f(y) \rangle| \leq \frac{L\|x-y\|^2}{2}. \quad (133)$$

To obtain the third inequality, we use Lemma C.2 to show that,

$$\|\nabla f(x)\| \leq G + L\|x\| \leq LR + G, \quad (134)$$

where the last inequality is due to fact that the domain is a ball with radius R . Then,

$$\|\nabla f(x)\|^2 - \|\nabla f(y)\|^2 = \|\nabla f(x) - \nabla f(y)\| \|\nabla f(x) + \nabla f(y)\| \leq 2(LR+G)L\|x-y\|. \quad (135)$$

Consequently, $e(x, y) \leq \frac{\beta L\|x-y\|^2}{2} + \frac{\beta\eta L(LR+G)\|x-y\|}{2}$. Finally, we need to bound

$$\int (\sqrt{p_{xy}^*} - \sqrt{p_{xy}})^2 dy \leq \frac{\mathbb{E}_{q(x,\cdot)}(e(x, y)^2)}{2} + \frac{\mathbb{E}_{q(x,\cdot)}(e(x, y))^2}{8} \quad (136)$$

$$\leq \frac{5}{8} \mathbb{E}_{q_{xy}}(e(x, y)^2) \quad (137)$$

$$\leq \frac{5}{8} \mathbb{E}_{q_{xy}} \left(\frac{\beta L\|x-y\|^2}{2} + \frac{\beta\eta L(LR+G)\|x-y\|}{2} \right)^2 \quad (138)$$

$$\leq \frac{5}{8} \beta^2 L^2 \mathbb{E}_{q_{xy}} \|x-y\|^4 + \frac{5}{8} \beta^2 \eta^2 L^2 (LR+G)^2 \mathbb{E}_{q_{xy}} \|x-y\|^2, \quad (139)$$

where the second inequality uses Jensen's inequality due to convexity of $e(x, y)$ and the last inequality is due to Young's inequality. Next, we need to compute the expectation values. Notice that since q_{xy} is a Gaussian, the variable $\frac{\beta\|y-x+\nabla f(x)\|^2}{\eta}$ is a chi-squared distributed random variable with mean d and variance $2d$.

$$\mathbb{E}_{q_{xy}} \|x-y\|^2 = \mathbb{E}_{q_{xy}} \|x-y + \eta\nabla f(x) - \eta\nabla f(x)\|^2 \quad (140)$$

$$\leq 2\mathbb{E}_{q_{xy}} \|x-y - \eta\nabla f(x)\|^2 + 2\eta^2 \mathbb{E}_{q_{xy}} \|\nabla f(x)\|^2 \quad (141)$$

$$\leq 2\eta d/\beta + 2\eta^2 (LR+G)^2 \quad (142)$$

$$\leq 4\eta d/\beta, \quad (143)$$

since the mean of chi squared distribution is d and $\eta \leq \frac{d}{\beta(LR+G)^2}$. Furthermore,

$$\mathbb{E}_{q_{xy}} \|x - y\|^4 = \text{Var}_{q_{xy}} \|x - y\|^2 + (\mathbb{E}_{q_{xy}} \|x - y\|^2)^2 \quad (144)$$

$$\leq 2d\eta^2/\beta^2 + (4\eta d/\beta)^2 \quad (145)$$

$$\leq 2d\eta^2/\beta^2 + 16\eta^2 d^2/\beta^2, \quad (146)$$

since variance of chi squared distribution is $2d$. Putting things together, we have for any $x \in \Omega$,

$$\|P(x, \cdot) - \tilde{P}(x, \cdot)\|_H^2 \leq \frac{5d\eta^2 L^2}{4} + 10\eta^2 d^2 L^2 + \frac{5\eta^3 d\beta L^2 (LR+G)^2}{2} \quad (147)$$

$$\leq 16\eta^2 L^2 d^2, \quad (148)$$

for $\eta \leq \frac{d}{\beta(LR+G)^2}$. Hence, $\|P(x, \cdot) - \tilde{P}(x, \cdot)\|_H \leq 4\eta dL$. $\square$

As the quantum walk operator is the basic building block of the reflection operators used in amplitude amplification, we present the following result to relate the error in quantum walk operator to the projection operators.

Lemma D.6. *Let W be a unitary operator with phase gap Δ and assume that W has a unique eigenvector $|\psi_0\rangle$ with eigenvalue 1. Suppose that we have $\tilde{W}$ such that,*

$$\|W - \tilde{W}\| \leq \delta. \quad (149)$$

Let $\Pi_{<\Delta}$ and $\tilde{\Pi}_{<\Delta}$ be operators that project any quantum state onto the space of eigenvectors of W and $\tilde{W}$ with phases smaller than Δ respectively. Then,

$$\|\Pi_{<\Delta} - \tilde{\Pi}_{<\Delta}\| \leq \frac{\delta\pi}{4\Delta}. \quad (150)$$

Proof. Let $W = \sum_m e^{2i\phi_m} |\psi_m\rangle \langle \psi_m|$ where $\phi_0 = 0$. Similarly, let $\tilde{W} = \sum_m e^{2i\tilde{\phi}_m} |\tilde{\psi}_m\rangle \langle \tilde{\psi}_m|$.

$$\|W\psi_0 - \tilde{W}\psi_0\|^2 = \left\| \sum_m \left(1 - e^{2i\tilde{\phi}_m}\right) |\tilde{\psi}_m\rangle \langle \tilde{\psi}_m|\psi_0\rangle \right\|^2 \quad (151)$$

$$= \sum_m |1 - e^{2i\tilde{\phi}_m}|^2 \left| \langle \tilde{\psi}_m|\psi_0\rangle \right|^2 \quad (152)$$

$$\geq \sum_{m:\tilde{\phi}_m \geq \Delta} |1 - e^{2i\tilde{\phi}_m}|^2 \left| \langle \tilde{\psi}_m|\psi_0\rangle \right|^2 \quad (153)$$

$$\geq 16\Delta^2/\pi^2 \sum_{m:\tilde{\phi}_m \geq \Delta} \left| \langle \tilde{\psi}_m|\psi_0\rangle \right|^2, \quad (154)$$

where the second inequality is due to $|1 - e^{ix}| \geq 2|x|/\pi$ whenever $-\pi \leq x \leq \pi$. Since $\|W - \tilde{W}\| \leq \delta$, we have

$$\sum_{m:\tilde{\phi}_m < \Delta} \left| \langle \tilde{\psi}_m|\psi_0\rangle \right|^2 \geq 1 - \frac{\delta^2 \pi^2}{16\Delta^2}. \quad (155)$$

Let $|\chi\rangle = \alpha_0 |\psi_0\rangle + \alpha_1 |\psi_0^\perp\rangle$ be an arbitrary quantum state such that $\alpha_1, \alpha_2 \in \mathbb{C}$ and $|\alpha_1|^2 + |\alpha_2|^2 = 1$. Then due to triangular inequality

$$\|\Pi_{<\Delta} |\chi\rangle - \tilde{\Pi}_{<\Delta} |\chi\rangle\| \leq |\alpha_0| \|\Pi_{<\Delta} |\psi_0\rangle - \tilde{\Pi}_{<\Delta} |\psi_0\rangle\| + |\alpha_1| \|\Pi_{<\Delta} |\psi_0^\perp\rangle - \tilde{\Pi}_{<\Delta} |\psi_0^\perp\rangle\|. \quad (156)$$

We first focus on the first term:

$$\|\Pi_{<\Delta} |\psi_0\rangle - \tilde{\Pi}_{<\Delta} |\psi_0\rangle\| = \left\| |\psi_0\rangle - \sum_{m:\tilde{\phi}_m < \Delta} |\tilde{\psi}_m\rangle \langle \tilde{\psi}_m|\psi_0\rangle \right\| \quad (157)$$

$$= \left(2 - 2 \sum_{m:\tilde{\phi}_m < \Delta} \left| \langle \tilde{\psi}_m|\psi_0\rangle \right|^2 \right)^{1/2} \quad (158)$$

$$\leq \frac{\delta\pi}{4\Delta}. \quad (159)$$

Similarly, for the second term,

$$\|\Pi_{<\Delta} |\psi_0^\perp\rangle - \tilde{\Pi}_{<\Delta} |\psi_0^\perp\rangle\| = \left\| \sum_{m: \tilde{\phi}_m < \Delta} |\tilde{\psi}_m\rangle \langle \tilde{\psi}_m | \psi_0^\perp \rangle \right\| \quad (160)$$

$$= \left(\sum_{m: \tilde{\phi}_m < \Delta} \left| \langle \tilde{\psi}_m | \psi_0^\perp \rangle \right|^2 \right)^{1/2} \quad (161)$$

$$= \left(1 - \sum_{m: \tilde{\phi}_m \geq \Delta} \left| \langle \tilde{\psi}_m | \psi_0 \rangle \right|^2 \right)^{1/2} \quad (162)$$

$$\leq \frac{\delta\pi}{4\Delta}. \quad (163)$$

Since both terms are smaller than $\frac{\delta\pi}{4\Delta}$, we conclude that for any state $|\chi\rangle$, the projectors are at most $\delta\pi/(4\Delta)$ apart in spectral norm. $\square$

Next lemma, also proved in Appendix D.2, quantifies the number of required controlled- U operators to implement the reflection operators.

Lemma D.7. *Let U be the quantum walk operator associated with Unadjusted Langevin algorithm. Under Assumptions 1.1 and 1.2 the reflection operator $V = e^{i\pi/3} |\pi\rangle \langle \pi| + (I - |\pi\rangle \langle \pi|)$ can be implemented with ϵ accuracy in spectral norm using $\tilde{O}(\rho^{-1}\beta dL\epsilon^{-1})$ controlled- U operators.*

Proof. Let P^* and P be the transition density of Metropolis Adjusted Langevin algorithm and Unadjusted Langevin algorithm respectively. Let U^* and U be the quantum walk operators built using P^* and P respectively. We can write U^* in spectral form:

$$U^* = \sum_m e^{2i\phi_m} |\psi_m\rangle \langle \psi_m|. \quad (164)$$

The phase gap Δ of U^* is defined to be $2|\phi_1|$. Since P^* is a reversible Markov chain, U^* accepts $|\pi\rangle$ as its eigenvector with eigenvalue 1. Furthermore, $|\pi\rangle$ is the unique eigenvector of U^* with eigenvalue 1 (see (Magniez et al., 2011) for more details). Notice that, R can be written as,

$$V = e^{i\pi/3} \Pi_\Delta^* + (I - \Pi_\Delta^*), \quad (165)$$

where Π_Δ^* is the projector that projects any quantum state onto the eigenstate of U^* with eigenphase smaller than Δ . This is because the only eigenvector of U^* with phase smaller than Δ is $|\pi\rangle$. This operator can be implemented in ϵ accuracy using techniques such as quantum singular value transformation technique introduced in (Gilyén et al., 2019) or phase estimation based method ((Magniez et al., 2011)) using $\tilde{O}(1/\Delta)$ calls to quantum walk operator. Suppose that we replaced each U^* with U and implement the following operator instead:

$$\tilde{V} = e^{i\pi/3} \Pi_\Delta + (I - \Pi_\Delta), \quad (166)$$

where Π is the projector similarly defined for U which is the quantum walk operator constructed for the unadjusted Langevin algorithm. Therefore, we can characterize the error,

$$\|V - \tilde{V}\| \leq 2\|\Pi_\Delta^* - \Pi_\Delta\| \quad (167)$$

$$\leq \frac{\|U - U^*\|}{2\Delta/\pi}. \quad (168)$$

The last inequality follows from Lemma D.6. By Lemma D.3 and Lemma D.2, $\Delta(U^*) \geq c_0\rho\sqrt{2\eta/\beta}$ for step size smaller than $O(d^{-1}\beta^{-1})$. Therefore,

$$\|V - \tilde{V}\| \leq c_0 16\sqrt{2\pi\eta} dL/\Delta \quad (169)$$

$$\leq c_0 16\pi\sqrt{2\eta\beta} dL/\rho, \quad (170)$$

where the first inequality is due to Lemma D.4 and Lemma D.5. Therefore, by setting $\eta \leq \frac{\epsilon^2 \rho^2}{c_0 16 \sqrt{2} \pi d^2 L^2 \beta}$, we have

$$\|V - \tilde{V}\| \leq \epsilon. \quad (171)$$

The total number of calls to U is $\tilde{O}(1/\Delta) = \tilde{O}(\rho^{-1} \eta^{-1/2} / \beta^{-1/2}) = \tilde{O}(\rho^{-1} \beta d L / \epsilon)$. $\square$

We are now ready to prove the query complexity of quantum ULA algorithm. We restate our result and give its proof next.

Theorem 4.2 (Quantum ULA). *Let $\pi \propto e^{-\beta f(x)}$ denote a probability distribution with inverse temperature $\beta > 0$ such that $f(x)$ satisfies Assumptions 1.1 and 1.2. Then, there exists a quantum algorithm that outputs a random variable distributed according to μ such that,*

$$\|\mu - \pi\|_{\text{TV}} \leq \epsilon, \quad (27)$$

where $\|\cdot\|_{\text{TV}}$ is the total variation distance, using $\tilde{O}(\beta d^{3/2} \epsilon^{-1} \rho^{-1} c_{\text{LSI}}^{-1})$ queries to $\mathcal{O}_{\nabla f}$.

Proof. Let $P^*(x \rightarrow y) = p_{xy}^*$ and $P(x \rightarrow y) = p_{xy}$ denote the transition densities of MALA and ULA algorithms respectively. Similarly, let U^* and U be the quantum walk operators associated with P^* and P constructed. We use the same algorithm described in proof of quantum MALA algorithm. That is, we iteratively drive each state $|\mu_i\rangle$ to $|\mu_{i+1}\rangle$ using $\pi/3$ fixed point amplitude amplification algorithm. However, since accessing U^* requires evaluation oracle, we instead use U to implement the reflection operator inexactly. The reflection operators can be implemented using $\tilde{O}(\rho^{-1} \beta d L \epsilon^{-1})$ calls to controlled U operator by Lemma D.7. Since the length of annealing schedule in Lemma 3.1 is $\tilde{O}(c_{\text{LSI}}^{-1} \sqrt{d})$, the total complexity is $\tilde{O}(c_{\text{LSI}}^{-1} \rho^{-1} d^{3/2} \beta \epsilon^{-1})$. Implementing U only requires full gradient oracle constant number of times, we establish the result. $\square$

D.3. Proofs for Quantum Stochastic ULA

The next lemma, proved in Appendix D.3, quantifies the expectation value of U_ℓ over ℓ with respect to a deterministic unitary U .

Lemma D.8. *Let $U_\ell = S\left(2 \sum_x |\psi_x^{(\ell)}\rangle \langle \psi_x^{(\ell)}| - I\right)$ be a quantum walk operator where $|\psi_x^{(\ell)}\rangle = \sum_y \sqrt{p_{xy}^{(\ell)}} |y\rangle$ is a quantum state constructed with stochastic gradient g_ℓ . Let $U = S\left(2 \sum_x |\psi_x\rangle \langle \psi_x| - I\right)$. Then, we have*

$$\|\mathbb{E}_\ell U_\ell - U\| \leq 6 \max_{x \in \Omega} \|\mathbb{E}_\ell |\psi_x^{(\ell)}\rangle - |\psi_x\rangle\|. \quad (172)$$

Proof.

$$\|\mathbb{E}_\ell U_\ell - U\| \leq 2 \left\| \mathbb{E}_\ell \sum_x |\psi_x^{(\ell)}\rangle \langle \psi_x^{(\ell)}| - \sum_x |\psi_x\rangle \langle \psi_x| \right\| \quad (173)$$

$$= 2 \left\| \mathbb{E}_\ell \sum_x |\psi_x^{(\ell)}\rangle (\langle \psi_x^{(\ell)}| - \langle \psi_x|) + \sum_x (|\psi_x^{(\ell)}\rangle - |\psi_x\rangle) \langle \psi_x| \right\| \quad (174)$$

$$\leq 2 \left\| \mathbb{E}_\ell \sum_x |\psi_x^{(\ell)}\rangle (\langle \psi_x^{(\ell)}| - \langle \psi_x|) \right\| + 2 \left\| \sum_x (|\psi_x^{(\ell)}\rangle - |\psi_x\rangle) \langle \psi_x| \right\|, \quad (175)$$

where the second inequality follows from triangular inequality. First, we focus on the second term,

$$\left\| \mathbb{E}_\ell \sum_x (|\psi_x^{(\ell)}\rangle - |\psi_x\rangle) \langle \psi_x| \right\| = \max_{|\phi\rangle} \left\| \mathbb{E}_\ell \sum_x (|\psi_x^{(\ell)}\rangle - |\psi_x\rangle) \langle \psi_x | \phi \rangle \right\|. \quad (176)$$

We can expand the state that maximizes this equation as $|\phi\rangle = \sum_x c_x |\psi_x\rangle + |\xi\rangle$ where $\langle \xi | \psi_x \rangle = 0$ for any $x \in \Omega$. This is true because $\langle \psi_x | \psi_y \rangle = \delta_{xy}$. Therefore, $\langle \psi_x | \phi \rangle = c_x$. Then,

$$\left\| \mathbb{E}_\ell \sum_x (|\psi_x^{(\ell)}\rangle - |\psi_x\rangle) \langle \psi_x| \right\| = \left\| \mathbb{E}_\ell \sum_x c_x (|\psi_x^{(\ell)}\rangle - |\psi_x\rangle) \right\| \quad (177)$$

$$= \left(\sum_x |c_x|^2 \|\mathbb{E}_\ell |\psi_x^{(\ell)}\rangle - |\psi_x\rangle\|^2 \right)^{1/2} \quad (178)$$

$$\leq \max_x \left(\|\mathbb{E}_\ell |\psi_x^{(\ell)}\rangle - |\psi_x\rangle\|^2 \right)^{1/2} \quad (179)$$

$$= \max_x \|\mathbb{E}_\ell |\psi_x^{(\ell)}\rangle - |\psi_x\rangle\|. \quad (180)$$

Again, the first equality is due to $\langle \psi_x | \psi_y \rangle = \delta_{xy}$ and the first inequality is due to fact that $\sum_x |c_x|^2 \leq 1$. The first term can be written as

$$2 \left\| \mathbb{E}_\ell \sum_x |\psi_x^{(\ell)}\rangle (\langle \psi_x^{(\ell)}| - \langle \psi_x|) \right\| = 2 \left\| \mathbb{E}_\ell \sum_x (|\psi_x^{(\ell)}\rangle - |\psi_x\rangle) (\langle \psi_x^{(\ell)}| - \langle \psi_x|) + \mathbb{E}_\ell \sum_x |\psi_x\rangle (\langle \psi_x^{(\ell)}| - \langle \psi_x|) \right\| \quad (181)$$

$$\leq 2 \left\| \mathbb{E}_\ell \sum_x (|\psi_x^{(\ell)}\rangle - |\psi_x\rangle) (\langle \psi_x^{(\ell)}| - \langle \psi_x|) \right\| \quad (182)$$

$$+ 2 \left\| \mathbb{E}_\ell \sum_x |\psi_x\rangle (\langle \psi_x^{(\ell)}| - \langle \psi_x|) \right\|. \quad (183)$$

The second term is bounded by $\max_x \|\psi_x - \mathbb{E}_\ell \psi_x^{(\ell)}\|$ and the first term,

$$\left\| \mathbb{E}_\ell \sum_x (|\psi_x\rangle - |\psi_x^{(\ell)}\rangle) (\langle \psi_x| - \langle \psi_x^{(\ell)}|) \right\| \leq \max_x \left\| \mathbb{E}_\ell (|\psi_x\rangle - |\psi_x^{(\ell)}\rangle) (\langle \psi_x| - \langle \psi_x^{(\ell)}|) \right\| \quad (184)$$

$$= \max_x \max_{|\phi\rangle} \|\mathbb{E}_\ell (|\psi_x\rangle - |\psi_x^{(\ell)}\rangle) (\langle \psi_x| - \langle \psi_x^{(\ell)}|) |\phi\rangle\| \quad (185)$$

$$\leq \max_x \|\mathbb{E}_\ell (|\psi_x\rangle - |\psi_x^{(\ell)}\rangle)\|, \quad (186)$$

the first inequality is due to the fact that for different $x, y \in \Omega$, $(\langle \psi_x| - \langle \psi_x^{(\ell)}|)(|\psi_y\rangle - |\psi_y^{(\ell)}\rangle) = 0$ and the last inequality is because $|\langle \psi_x| - \langle \psi_x^{(\ell)}| |\phi\rangle| \leq 1$. □

The next lemma is the application of Lemma D.8 on quantum Langevin algorithms.

Lemma D.9. *Let U be the quantum walk operator for unadjusted Langevin algorithm computed using exact gradients. Let U_ℓ be a quantum walk operator for unadjusted Langevin algorithm constructed by computing the gradient on random mini batch ℓ of size B . Then, under Assumptions 1.1 and 1.2, we have*

$$\|\mathbb{E}_\ell U_\ell - U\| \leq 6\sqrt{2}\eta\beta(LR + G)d^{1/2}/B^{1/2}, \quad (187)$$

where G is a positive constant that satisfies $\|\nabla f(0)\| \leq G$.

Proof.

$$\|U - \mathbb{E}_\ell U_\ell\|^2 \leq 36 \max_{x \in \Omega} \|\psi_x - \mathbb{E}_\ell \psi_x^{(\ell)}\|^2 \quad (188)$$

$$\leq 36 \max_{x \in \Omega} \left\| \int_{y \in \mathbb{R}^d} dy (\sqrt{p_{xy}} - \mathbb{E}_\ell \sqrt{p_{xy}^\ell}) |x\rangle |y\rangle \right\|^2 \quad (189)$$

$$= 36 \max_{x \in \Omega} \left(\int_{y \in \mathbb{R}^d} dy p_{xy} + \int_{y \in \mathbb{R}^d} dy (\mathbb{E}_\ell \sqrt{p_{xy}^\ell})^2 - 2\mathbb{E}_\ell \int_{y \in \mathbb{R}^d} dy \sqrt{p_{xy} p_{xy}^\ell} \right) \quad (190)$$

$$\leq 36 \max_{x \in \Omega} \left(\int_{y \in \mathbb{R}^d} \sqrt{p_{xy}} + \mathbb{E}_\ell \int_{y \in \mathbb{R}^d} dy p_{xy}^\ell - 2\mathbb{E}_\ell \int_{y \in \mathbb{R}^d} dy \sqrt{p_{xy} p_{xy}^\ell} \right) \quad (191)$$

$$= 36 \max_{x \in \Omega} \left(2 - 2\mathbb{E}_\ell \int_{y \in \mathbb{R}^d} dy \sqrt{p_{xy} p_{xy}^\ell} \right), \quad (192)$$

where the first inequality is due to [Lemma D.8](#) and the second inequality is due to Jensen's inequality since square root is a concave function.

$$\int_{y \in \mathbb{R}^d} dy \sqrt{p_{xy} p_{xy}^\ell} = \frac{1}{(4\pi\eta/\beta)^{d/2}} \int_{y \in \mathbb{R}^d} dy \exp\left(-\frac{\|y - x + \eta \nabla f(x)\|^2}{4\eta/\beta}\right) \exp\left(-\frac{\|y - x + \eta g_\ell(x)\|^2}{4\eta/\beta}\right) \quad (193)$$

$$= \frac{1}{(4\pi\eta/\beta)^{d/2}} \int_{y \in \mathbb{R}^d} dy \exp\left(-\frac{2\|y - x\|^2 + 2\eta \langle y - x, \nabla f(x) + g_\ell(x) \rangle + \eta^2 \|\nabla f(x)\|^2 + \eta^2 \|g_\ell(x)\|^2}{4\eta/\beta}\right) \quad (194)$$

$$= \frac{1}{(4\pi\eta/\beta)^{d/2}} \int_{y \in \mathbb{R}^d} dy \exp\left(-\frac{\|y - x + \eta(\nabla f(x) + g_\ell(x))/2\|^2}{2\eta/\beta}\right) \exp\left(-\frac{\eta^2 \|\nabla f(x) - g_\ell(x)\|^2}{2\eta/\beta}\right) \quad (195)$$

$$= \exp\left(-\frac{\eta^2 \|\nabla f(x) - g_\ell(x)\|^2}{2\eta/\beta}\right), \quad (196)$$

where $\mathbb{E}[g_\ell] = \nabla f$. Therefore,

$$\|U - \mathbb{E}U_\ell\|^2 \leq 36 \max_{x \in \Omega} \left(2 - 2\mathbb{E} \exp\left(-\frac{\eta^2 \|\nabla f(x) - g_\ell(x)\|^2}{2\eta/\beta}\right)\right) \quad (197)$$

$$\leq 36(2 - 2\exp(-\eta^2 \beta^2 (LR + G)^2/B)) \quad (198)$$

$$\leq 72\eta^2 d \beta^2 (LR + G)^2/B, \quad (199)$$

where the first inequality follows from lemma B.2 from [\(Zou et al., 2021\)](#),

$$\mathbb{E} \exp(\langle a, g_\ell(x) - \nabla f \rangle) \leq \exp(M^2 \|a\|_2^2/B), \quad (200)$$

where M is the upper bound on $\|g_\ell(x) - \nabla f(x)\|$ with batch size B . $\square$

The next lemma upper bounds the difference of two random unitary quantum walk operators.

Lemma D.10. *Let U_{ℓ_1} and U_{ℓ_2} be two random quantum walk operators constructed with two different stochastic gradients g_{ℓ_1} and g_{ℓ_2} for unadjusted Langevin algorithm. Then, under Assumptions 1.1 and 1.2, we have*

$$\|U_{\ell_1} - U_{\ell_2}\| \leq 8\sqrt{\eta\beta(LR + G)^2}, \quad (201)$$

where G is a positive constant that satisfies $\|\nabla f(0)\| \leq G$.

Proof. By Lemma D.4, the difference of quantum walk operators is bounded by,

$$\|U_{\ell_1} - U_{\ell_2}\|^2 \leq 32 \max_x \|P_{\ell_1} - P_{\ell_2}\|_H^2, \quad (202)$$

where P_{ℓ_1} and P_{ℓ_2} are Gaussian transition densities of ULA computed with gradients on mini batches ℓ_1 and ℓ_2 . This is squared Hellinger distance between two Gaussian distributions with the same variance and different mean. This is a known result [\(Pardo, 2018\)](#) and equal to following.

$$\|P_{\ell_1} - P_{\ell_2}\|_H = 1 - \exp\left(-\frac{\eta^2 \|g_{\ell_1}(x) - g_{\ell_2}(x)\|^2}{2\eta/\beta}\right) \leq \frac{\eta^2 \|g_{\ell_1}(x) - g_{\ell_2}(x)\|^2}{2\eta/\beta}. \quad (203)$$

Since $\|\nabla f(x)\| \leq L\|x\| + G \leq LR + G$, $\|g_{\ell_1}(x) - g_{\ell_2}(x)\| \leq 2(LR + G)$, therefore for any $x \in \Omega$,

$$\|U_{\ell_1} - U_{\ell_2}\|^2 \leq 64(LR + G)^2 \eta \beta. \quad (204)$$

Taking the square root, we obtain the result in the statement. $\square$

Finally we prove the following theorem to conclude the analysis of stochastic quantum sampling algorithm.

Theorem 4.3 (Quantum ULA with stochastic gradient). *Let $\pi \propto e^{-\beta f(x)}$ denote a probability distribution with inverse temperature $\beta > 0$ such that $f(x) = \frac{1}{N} \sum_{k=1}^N f_k(x)$ satisfies Assumptions 1.1 and 1.2. Then, there exists a quantum algorithm that outputs a random variable distributed according to μ such that,*

$$\|\mu - \pi\|_{\text{TV}} \leq \epsilon, \quad (28)$$

where $\|\cdot\|_{\text{TV}}$ is the total variation distance, using $\tilde{O}(\beta^2 d^{3/2} \epsilon^{-2} \rho^{-2} c_{\text{LSI}}^{-1})^3$ queries to $\mathcal{O}_{\tilde{\nabla}f}$ and each $\mathcal{O}_{\tilde{\nabla}f}$ involves $O(d)$ gradient calculations.

Proof. Let U_ℓ be a unitary quantum walk operator defined as,

$$U_\ell = S \left(2 \sum_x |\psi_x^{(\ell)}\rangle \langle \psi_x^{(\ell)}| - I \right), \quad (205)$$

where $|\psi_x^{(\ell)}\rangle$ is the state,

$$|\psi_x^{(\ell)}\rangle = \sum_y \sqrt{p_{xy}^{(\ell)}} |x\rangle |y\rangle, \quad (206)$$

where $p_{xy}^{(\ell)} = \frac{1}{(4\pi\eta/\beta)} \exp\left(-\frac{\|y-x+g_\ell(x)\|^2}{2\eta/\beta}\right)$, and g_ℓ is the stochastic gradient computed on randomly selected data points of size B , i.e.,

$$g_\ell(x) := \frac{1}{B} \sum_{i \in S_\ell \subseteq [N]} \nabla f_i(x). \quad (207)$$

The number of gradient evaluations for implementing unitary U_ℓ is $O(B)$ since we only need to compute gradient on B data points. The key idea in proof of the quantum ULA is the fact that the following operator can be implemented using controlled- U operators:

$$V = e^{i\pi/3} \Pi_\Delta + \Pi_\Delta^\perp, \quad (208)$$

where Δ is the phase gap of quantum MALA walk operator. Suppose that we replace every controlled- U operator with a unitary U_ℓ . Note that each U in the circuit might be possibly replaced by different unitary due to randomness of stochastic gradients. Let's denote this circuit by $\tilde{V}$. Now, we show that with high probability $\|V - \tilde{V}\| \leq \epsilon$ for sufficiently small step size. Since the algorithm uses $1/\Delta(U^*)$ calls to U ,

$$\|V - \mathbb{E}(\tilde{V})\| \leq \frac{1}{\Delta} \|U^* - \mathbb{E}_\ell U_\ell\| \quad (209)$$

$$\leq \frac{1}{\Delta} \|U - U^*\| + \|U - \mathbb{E}_\ell U_\ell\| \quad (210)$$

$$\leq (\rho^{-1} \sqrt{\beta/\eta}) \eta d L + (\rho^{-1} \sqrt{\beta/\eta}) (\eta \beta (LR + G) \sqrt{d/B}) \quad (211)$$

$$= \rho^{-1} \eta^{1/2} \beta^{1/2} d L + \rho^{-1} \beta^{3/2} \eta^{1/2} (LR + G) d^{1/2} / B^{1/2}. \quad (212)$$

Setting $\eta \leq \min\left(\frac{\epsilon^2 \rho^2}{2\beta d^2 L^2}, \frac{\epsilon^4 \rho^2 B}{4\beta^3 d (LR+G)^2}\right)$ and $B = d$, we guarantee that,

$$\|V - \mathbb{E}\tilde{V}\| \leq \epsilon/2. \quad (213)$$

³As each $\mathcal{O}_{\tilde{\nabla}f}$ uses $\tilde{O}(d)$ gradient calculations, the number of total gradient calculations scale as $d^{5/2}$ as shown Table 1.

Next, we use the McDiarmid's inequality to obtain high probability bound:

$$P(\|\tilde{V} - \mathbb{E}\tilde{V}\| \geq \epsilon/2) \leq 2 \exp\left(-\frac{\epsilon^2 \Delta}{2\|U_{\ell_1} - U_{\ell_2}\|^2}\right) \quad (214)$$

$$\leq 2 \exp\left(-\frac{\epsilon^2 \rho \eta^{1/2}}{2\beta^{1/2}\|U_{\ell_1} - U_{\ell_2}\|^2}\right) \quad (215)$$

$$\leq 2 \exp\left(-\frac{\epsilon^2 \rho \eta^{1/2}}{128\beta^{1/2}\eta\beta(LR+G)^2}\right) \quad (216)$$

$$= 2 \exp\left(-\frac{\epsilon^2 \rho}{128\beta^{3/2}\eta^{1/2}(LR+G)^2}\right). \quad (217)$$

Setting $\eta \leq \frac{\epsilon^4 \rho^2}{128^2(LR+G)^4 \beta^3}$, guarantees that with at least constant probability,

$$\|\tilde{V} - \mathbb{E}\tilde{V}\| \leq \epsilon/2. \quad (218)$$

The probability can be boosted in logarithmic number of steps to obtain high probability. Therefore, with high probability,

$$\|V - \tilde{V}\| \leq \|V - \mathbb{E}\tilde{V}\| + \|\mathbb{E}\tilde{V} - \tilde{V}\| \leq \epsilon. \quad (219)$$

Then, to implement the operator V up to ϵ accuracy with high probability, we need $\tilde{O}(1/\Delta) = (\rho^{-1}\eta^{1/2}\beta^{-1/2}) = \tilde{O}(\rho^{-2}d\beta/\epsilon^2)$ calls to U_ℓ . Since each U_ℓ requires $B = d$ gradient computations and we need to prepare $\tilde{O}(c_{\text{LSI}}^{-1}\sqrt{d})$ reflections, the total gradient complexity is $\tilde{O}(c_{\text{LSI}}^{-1}\rho^{-2}d^{5/2}/\epsilon^2)$. $\square$

E. Proofs for Partition Function Estimation

Before we prove the main theorem for partition functions, we give the following lemmas. The following lemma shows that Z_1 can be estimated up to ϵ multiplicative constant from a Gaussian.

Lemma E.1 (Lemma 3.1 of (Ge et al., 2020)). *Letting $\sigma_1^2 = \frac{\epsilon}{2dL}$, it holds that*

$$\left(1 - \frac{\epsilon}{2}\right) \int_{x \in \mathbb{R}^d} \exp\left(-\frac{\|x\|^2}{2\sigma_1^2}\right) dx \leq Z_1 \leq \int_{x \in \mathbb{R}^d} \exp\left(-\frac{\|x\|^2}{2\sigma_1^2}\right) dx. \quad (220)$$

The next lemma uses unbiased quantum mean estimation to compute the product of ℓ random variables.

Lemma E.2 (Theorem 3.3 of (Cornelissen & Hamoudi, 2023)). *Let $B > 1$ and $\epsilon \in (0, 1)$. Consider a sequence $X_1, \dots, X_\ell$ of ℓ independent random variables with support size n , bounded relative second moment $\frac{\mathbb{E}[X_i^2]}{\mathbb{E}[X_i]^2} \leq B$ and bounded fidelity $|\langle \pi_{X_i} | \pi_{X_{i+1}} \rangle|^2 \geq 1/B$ for all i . Denote their product as $X = X_1 \dots X_\ell$. Then, there exists a quantum algorithm that outputs a multiplicative-error estimate $\tilde{p}$ such that*

$$\left| \tilde{p} - \mathbb{E}\left[\prod_{i=1}^{\ell} X_i\right] \right| \leq \epsilon \mathbb{E}\left[\prod_{i=1}^{\ell} X_i\right] \quad (221)$$

with probability at least $2/3$. It uses $O(B)$ copies of $|\pi_{X_1}\rangle$ and $\tilde{O}(B^2\ell^{3/2}/\epsilon + B\ell \log(n))$ reflections through the states $|\mu_{X_1}\rangle, \dots, |\mu_{X_\ell}\rangle$ in expectation.

Finally, we combine our sampling algorithms with the annealing schedule and product estimator to obtain our result for the partition function estimation.

Theorem 5.1. *Let $Z = \int_x e^{-f(x)} dx$ be the partition with $f(x)$ function satisfying assumptions Assumptions 1.1 and 1.2. Then, there exists quantum algorithms that output an estimate $\tilde{Z}$ such that,*

$$(1 - \epsilon)Z \leq \tilde{Z} \leq (1 + \epsilon)Z \quad (32)$$

with probability at least $3/4$ using,

- $\tilde{O}(d^{5/4}\epsilon^{-1}\rho^{-1}c_{\text{LSI}}^{-1})$ queries to $\mathcal{O}_{\nabla f}$ and $\mathcal{O}_f$, or
- $\tilde{O}(d^{7/4}\epsilon^{-2}\rho^{-1}c_{\text{LSI}}^{-1})$ queries to $\mathcal{O}_{\nabla f}$, or
- $\tilde{O}(d^{11/4}\epsilon^{-3}\rho^{-2}c_{\text{LSI}}^{-1})$ queries to $\mathcal{O}_{\tilde{\nabla} f}$.

Proof. By Lemma E.1, we can estimate Z_1 with ϵ accuracy using normalization constant of Gaussian distribution with variance σ_1^2 . Then, we show that g_i has constant relative variance for all $i \in [M]$. Since the partition function can be written in telescoping product given in Equation (30), we can use Lemma E.2 to estimate the remaining product up to ϵ multiplicative constant with high probability. First, for g_M ,

$$\frac{\mathbb{E}_{\mu_M}[g_M^2]}{\mathbb{E}_{\mu_M}[g_M]^2} = \mathbb{E}_\pi \left[\exp\left(-\frac{1}{2\sigma_M^2}\|x\|^2\right) \right] \mathbb{E}_\pi \left[\exp\left(\frac{1}{2\sigma_M^2}\|x\|^2\right) \right] \quad (222)$$

$$\leq O(\exp(dL/(m\sigma_M^4 c_{\text{LSI}}^2))) \quad (223)$$

by Lemma C.3. Setting $\sigma_M^2 = \Omega(\sqrt{dL/(mc_{\text{LSI}}^2)})$ implies $\frac{\mathbb{E}_{\mu_M}[g_M^2]}{\mathbb{E}_{\mu_M}[g_M]^2} \leq O(1)$. Similarly, for $i \in [1, M-1]$,

$$\frac{\mathbb{E}_{\mu_i}[g_i^2]}{\mathbb{E}_{\mu_i}[g_i]^2} = \frac{\mathbb{E}_\pi \left[\exp\left(-\frac{(1+\alpha)}{2}\|x\|^2\right) \right] \mathbb{E}_\pi \left[\exp\left(-\frac{(1-\alpha)}{2}\|x\|^2\right) \right]}{(\mathbb{E}_\pi[\exp(-\|x\|^2/2)])^2} \quad (224)$$

$$\leq O(\exp(dL\alpha^2/(mc_{\text{LSI}}^2))) \quad (225)$$

by Lemma C.4. Therefore, for $\alpha^2 = \tilde{O}(mc_{\text{LSI}}^2/(dL))$, $\frac{\mathbb{E}_{\mu_i}[g_i^2]}{\mathbb{E}_{\mu_i}[g_i]^2} \leq O(1)$. Having established that the relative variance is constant for all g_i , by Lemma E.2, we conclude that the product form can be estimated by using $\tilde{O}(M^{3/2}/\epsilon) = \tilde{O}(d^{3/4}/\epsilon)$ reflection operators. We can choose quantum MALA, quantum ULA or stochastic quantum ULA algorithms to implement the reflection operators. Since the complexities given in Theorem 4.1, Theorem 4.2, Theorem 4.3 are the complexities of implementing reflection operator times M , we just need to multiply these results with $\tilde{O}(M^{1/2}) = \tilde{O}(d^{1/4})$ to conclude the proof. $\square$

F. Dependence on Isoperimetric Constants

The dependency on the isoperimetric constants for quantum MALA, ULA and stochastic ULA are $\rho^{-1}c_{\text{LSI}}^{-1}$, $\rho^{-1}c_{\text{LSI}}^{-1}$ and $\rho^{-2}c_{\text{LSI}}^{-1}$ respectively as given in the Theorems 4.1 to 4.3. On the other hand, the dependency for the classical algorithms in Table 1 are c_{LSI}^{-2} , c_{LSI}^{-2} and ρ^{-4} . Unfortunately, there is no tight relation between c_{LSI} and ρ and this makes hard to make comparison without further structure on f . Although it is not fully rigorous, it is still possible to make a comparison by converting both Cheeger constants and log-Sobolev constants to Poincare constant (c_p) which is another way of expressing the global properties of the function landscape. Using $\rho \geq \Omega(d^{-1/2}c_p)$, and $c_{\text{LSI}} \geq c_p$ (Buser, 1982), we can show that quantum algorithms have the complexity $d^{3/2}c_p^{-2}$, $d^{5/2}c_p^{-2}$, $d^{7/2}c_p^{-3}$ for quantum MALA, ULA and stochastic ULA. On the other hand, the classical complexities become $d^2c_p^{-2}$, dc_p^{-2} and $d^6c_p^{-4}$. Note that this conversion does not change ϵ dependency. Hence, our quantum algorithm have the same dependency on c_p for MALA and ULA, whereas it has better dependency (c_p^{-3} vs c_p^{-4}) for stochastic ULA. As claimed in the main text, by expressing the bounds in terms of c_p , we also maintain the improvement in d for MALA and stochastic ULA. Unfortunately, Buser's inequality is not always tight and this argument both loosens classical and quantum bounds. For the sake of keeping the bounds sharp, we did not include this kind of comparison in Table 1.

We also note that in the most general case, the isoperimetric constants may be exponentially small on d , L and b . Therefore, in stochastic case, we might obtain a significant speedup. However, since the runtime is dominated by c_{LSI}^{-1} and ρ^{-1} , the dimension speedup for quantum MALA may not be as important. However, for certain non-convex functions encountered in machine learning, the dependency on d might not be exponential. For example, for locally non-convex function, which models the Gaussian mixtures, considered in (Ma et al., 2019b), c_{LSI}^{-1} scales as $O(\exp(LR^2))$ where R is the radius of the non-convex region.