

Graph-Based Ascent Algorithms for Function Maximization

Muni Sreenivas Pydi¹, Varun Jog² and Po-Ling Loh³

Department of Electrical & Computer Engineering

University of Wisconsin - Madison

Madison, WI 53706

Email: ¹pydi@wisc.edu, ²vjog@ece.wisc.edu, ³loh@ece.wisc.edu

Abstract—We study the problem of finding the maximum of a function defined on the nodes of a connected graph. The goal is to identify a node where the function obtains its maximum. We focus on local iterative algorithms, which traverse the nodes of the graph along a path, and the next iterate is chosen from the neighbors of the current iterate with probability distribution determined by the function values at the current iterate and its neighbors. We study two algorithms corresponding to a Metropolis-Hastings random walk with different transition kernels: (i) The first algorithm is an exponentially weighted random walk governed by a parameter γ . (ii) The second algorithm is defined with respect to the graph Laplacian and a smoothness parameter k . We derive convergence rates for the two algorithms in terms of total variation distance and hitting times. We also provide simulations showing the relative convergence rates of our algorithms in comparison to an unbiased random walk, as a function of the smoothness of the graph function. Our algorithms may be categorized as a new class of “descent-based” methods for function maximization on the nodes of a graph.

I. INTRODUCTION

Social media and big data analysis have led to an explosion of graph-structured datasets in diverse domains such as neuroscience, economics, sensor networks, and databases. In many important instances, the object of interest may be expressed as a function defined on the vertices of the graph. For example, in social network analysis, it is useful to identify a node or nodes with a maximal centrality (or “influence”) score relative to the other nodes, for the purpose of performing strategic interventions [1], [2]. In electrical network modeling, the edges of a graph correspond to connections in an electrical grid, and the amount of power available at each node may similarly be represented using a graph-structured function [3]. This has led to a variety of statistical studies in graph function estimation, including Laplacian smoothing [4], [5] and graph trend filtering [6], when noisy measurements of graph characteristics may be available.

In this paper, we focus on the problem of obtaining the maximum value of a function defined on the nodes of a graph. We assume a setting where only local information is available, and any algorithm must make decisions based on graph information obtained by traversing the nodes of the graph in a sequential manner. For instance, such a setting may arise in web crawling [7], respondent-driven sampling [8], or other demographic studies [9]. Although basic properties of Markov chains imply that an unbiased

random walk on the nodes of a connected graph will eventually visit every node—hence produce a function maximum—we wish to use information about the value of the function on neighboring nodes to guide the path of the algorithm, borrowing ideas from the extensive literature on function optimization in the continuous domain.

For continuous optimization, several algorithms such as gradient descent, stochastic gradient descent, stochastic gradient Langevin dynamics, and momentum-based descent algorithms [10], [11], [12] provide guarantees for finding local and global optima when the function is sufficiently smooth. Common classes of smooth functions include Lipschitz functions, convex functions, and Polyak-Łojasiewicz functions [13]. Most of these notions, however, do not have direct analogs in the discrete domain of graphs. For instance, there is no obvious analog of a convex function on a graph [14], [15]. In this paper, we take a step toward systematically studying function optimization on graphs by considering a certain class of smooth graph functions, where the notion of smoothness is motivated by the theory of band-limited graph signals. Such topics constitute an area of study in graph signal-processing [16], which has become popular within the signal processing community in recent years [17], [18]. The algorithms we propose exploit the underlying graph structure to construct efficient algorithms for maximizing such smooth graph functions, producing a promising analogy to gradient descent algorithms for continuous functions. Two key features of all our algorithms are that they are iterative and local; i.e., the random walk relies only on the function values of neighboring nodes when deciding its next step. This is a direct analogy to first-order optimization methods in continuous domains.

Our proposed algorithms employ random walk methods based on Monte Carlo approaches, specifically the Metropolis-Hastings (MH) algorithm [19]. An important idea is to reformulate the problem of graph function maximization as a problem involving sampling vertices with high function values. Such strategies have recently gained traction in continuous domain optimization, as well [20], [21], [22]. Starting with a proposal probability density on the vertices of the graph that is related to the original function values in an appropriate manner, we construct random walk transition kernels that converge to the desired proposal distribution. Our approach differs significantly from comparable approaches in the literature in the following way: In constructing the MH

random walk kernel, we exploit the graph-connectivity information implied by the graph Laplacian matrix to construct rapidly-mixing walks that converge to the function maximum. Through hitting time analysis, we provide theoretical guarantees for the speed of convergence of these algorithms.

The remainder of the paper is organized as follows: In Section II, we introduce the mathematical formulation of the graph maximization problem and the types of local algorithms and smooth functions we will consider. In Section III, we provide a formal statement of our proposed algorithms, which are analyzed rigorously in Section IV. In Section V, we provide simulations demonstrating the behavior of our algorithms. We conclude with a list of interesting open questions in Section VI.

a) Notation: We write $\|\cdot\|_{TV}$ to denote the total variation norm and $\|\cdot\|_2$ to denote the Euclidean norm of a vector. For a matrix $A \in \mathbb{R}^{n \times n}$, we write $\text{diag}(A)$ to denote the $n \times n$ matrix with diagonal entries equal to the components of A and all other entries equal to zero. We write $\lfloor \cdot \rfloor$ to denote the floor operator.

II. BACKGROUND AND PROBLEM SETUP

Consider an unweighted, undirected graph G with vertex set V of size n and symmetric adjacency matrix W . For a node $i \in V$, we write d_i to denote the degree of i and $N(i)$ to denote the neighborhood set of i . We write $d_{\max} = \max_{i \in V} d_i$. We refer to a function $f : V \rightarrow \mathbb{R}$ as a *graph function* defined on the vertices of G , and sometimes write f_i to denote the function value $f(i)$. Our goal is to design a local algorithm that finds the maximum value of an unknown graph function in an efficient manner.

A. Local Search Algorithms

Formally, we define a local search algorithm as a discrete time Markov process, where the states of the process are the n nodes, and a Markov chain in state X_t at time t chooses the next state X_{t+1} among the neighbors of X_t , with transition probability distribution defined as a function of the value $f(X_t)$ and the set of values $\{f(V_t) : V_t \in N(X_t)\}$ at neighboring nodes in the graph.

We briefly mention some related work. Borgs et al. [23] studied the problem of finding the first node in a preferential attachment network, where the algorithm is given access to all nodes in a frontier of radius r around the set of previously queried nodes. Frieze and Pegden [24] proposed an alternative algorithm of finding the root of a preferential attachment graph, where the algorithm is given access to a sorted list of the highest-degree neighbors of the current iterate. Brautbar and Kearns [7] studied algorithms for finding nodes with high degrees or high clustering coefficients using a “jump and crawl” method, where a crawl step consists of randomly querying a neighbor of the existing iterate, and a jump step allows the algorithm to query a uniformly sampled node from the vertex set of the graph. Note, however, that our methods differ from the studies of Borgs et al. [23] and Frieze and Pegden [24] due to the fact that we are interested in different classes of graph functions besides the degree

function (indeed, the degree function may not be smooth), and from the study of Brautbar and Kearns [7] due to the fact that successive steps may only visit immediate neighbors, rather than jumping to an entirely different node in the graph.

B. Metropolis-Hastings Algorithm

Many of our algorithms are based on versions of the Metropolis-Hastings Algorithm, which we review in this section. Given a state space V and a target probability density $p_f(\cdot)$, the MH algorithm constructs a Markov chain with transition probability matrix $\mathbf{P}$ whose stationary distribution π is the same as the target probability density (i.e., $\pi(i) = p_f(i) \forall i \in V$). The matrix $\mathbf{P}$ is defined in terms of a “proposal” distribution $\mathbf{Q}$, as follows:

$$\mathbf{P}_{ij} = \begin{cases} \mathbf{Q}_{ij} R(i, j), & j \neq i \\ 1 - \sum_{j \neq i} \mathbf{P}_{ij}, & j = i, \end{cases} \quad (1)$$

where

$$R(i, j) = \min \left(1, \frac{p_f(j) \mathbf{Q}_{ji}}{p_f(i) \mathbf{Q}_{ij}} \right). \quad (2)$$

The general idea behind the MH algorithm is that successive steps are guided by the proposal distribution, such that the move from node i to node j depends on the ratio $\frac{p_f(j)}{p_f(i)}$, but is also moderated by the proposal distribution according to the ratio $\frac{\mathbf{Q}_{ji}}{\mathbf{Q}_{ij}}$. Known results in probability theory guarantee the convergence of the MH algorithm to the target density [25].

In addition to generating samples from a desired distribution, however, the MH algorithm is also a valuable tool in stochastic optimization [19]. The key idea is that if a function f with domain V is encoded into an appropriate density p_f , such that larger values of f correspond to larger values of p_f , the random walk generated by the MH algorithm will be more likely to visit nodes with larger density values, hence larger values of f .

The *independent MH algorithm* corresponds to the choice of proposal distribution $\mathbf{Q}_{ij} = p_g(j)$, where p_g is some probability measure over the state space V . In Section III, we present a variety of algorithms for function maximization based on running the MH algorithm with different choices of proposal distributions and target densities, for which we may obtain provable guarantees for the rate of convergence.

C. Smooth Graph Functions

We wish to leverage information about local changes in the graph function across edges in the graph to guide the movement of an algorithm. Intuitively, this will lead to improvements in the rate of convergence of an algorithm to the graph maximum when the function satisfies certain smoothness properties, as a rough analog of the convergence results of gradient descent for smooth functions on a real domain. In order to make these notions rigorous, we now define the smooth class of graph functions to be analyzed in our paper, borrowing from the literature on graph signal processing.

The unnormalized graph Laplacian for a graph G with adjacency matrix W is given by $L = D - W$, where $D =$

$\text{diag}\left(\sum_{j=1}^n W_{ij}\right)$ is the degree matrix. Since the Laplacian is symmetric, we can calculate the eigendecomposition of L as $U\Lambda U^T$, where $U \in \mathbb{R}^{n \times n}$ is an orthonormal matrix whose i^{th} column u_i is the eigenvector corresponding to the eigenvalue λ_i , and Λ is a diagonal matrix with sorted eigenvalues $0 = \lambda_1 \leq \dots \leq \lambda_n$.

Since U is unitary, the columns of U form an orthogonal basis for any graph function f . In the basis of Laplacian eigenvectors, the function f may be written as $f = U\hat{f}$, where $\hat{f}$ is defined as the graph Fourier transform $\hat{f} = U^T f$. The i^{th} column u_i of U is the i^{th} Fourier mode of the graph, corresponding to the i^{th} smallest graph frequency, also equal to $\sqrt{\lambda_i}$ [16].

Intuitively, a “smooth” graph function is a function that does not vary much across vertices that are connected by edges. This notion of smoothness is captured by the term $f^T L f$, as shown in the following identity [26]:

$$f^T L f = \frac{1}{2} \sum_{i,j=1}^n w_{ij} (f_i - f_j)^2. \quad (3)$$

For a smooth function, the sum of squared differences $(f_i - f_j)^2$ across all the connected vertices is small, implying that $f^T L f$ is small. Representing f in the Fourier basis, we have

$$f^T L f = (\hat{f}^T U^T) U \Lambda U^T (U \hat{f}) = \hat{f}^T \Lambda \hat{f} = \sum_{i=1}^n \lambda_i \hat{f}_i^2.$$

Since the λ_i 's are arranged in ascending order, setting the higher-order Fourier coefficients $\hat{f}_i$'s close to zero leads to a small value for $f^T L f$, implying that f is smooth. We have the following definition:

Definition 1. A function $f : G \rightarrow \mathbb{R}$ is called k -smooth if there exists a vector $\alpha \in \mathbb{R}^k$ such that $f = U_k \alpha$, where $U_k \in \mathbb{R}^{n \times k}$ is the matrix consisting of the first k columns of U .

In general, any graph function f can be decomposed as $f = f_{ks} + f_r$, where $f_{ks} = U_k U_k^T f$ is the k -smooth component and f_r is the residual component. Note that $f_r \perp f_{ks}$, so

$$\|f\|_2^2 = \|f_{ks}\|_2^2 + \|f_r\|_2^2.$$

In graph signal processing, a k -smooth function is viewed as a band-limited graph signal restricted to the lowest k Fourier modes of the graph. The vector α represents the projection of f onto the lowest k graph frequencies [18].

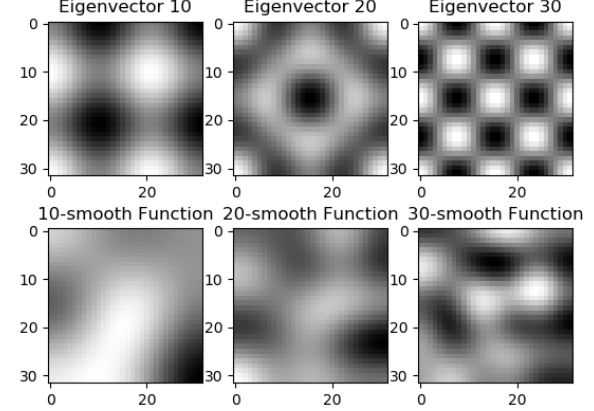
Denote δ_i to be the vector with 1 in the i^{th} coordinate and zeros elsewhere. We then have the following definition:

Definition 2. The term $\|U_k^T \delta_i\|_2$ is referred to as the local cumulative coherence of order k (LC- k) at node i [27]. The LC- k quantity satisfies

$$\|U_k^T \delta_i\|_2 = \frac{\|U_k^T \delta_i\|_2}{\|\delta_i\|_2} = \frac{\|U_k^T \delta_i\|_2}{\|U^T \delta_i\|_2}. \quad (4)$$

From equation (4), we see that $\|U_k^T \delta_i\|_2$ gives the proportion of the energy of the graph impulse function δ_i that is concentrated on the top k Fourier modes [18]. The LC- k

Fig. 1. Examples of randomly generated smooth functions with $k = 10$, $k = 20$, and $k = 30$ on a 32×32 2D grid graph of $n = 1024$ nodes. The eigenvectors u_{10} , u_{20} and u_{30} are plotted in the first row. A higher intensity on a pixel indicates a higher positive value for the vector at that node. Note that as expected, functions corresponding to a smaller value of k have fewer undulations.



varies between 0 and 1. Whereas we do not usually associate δ_i with smoothness, it may happen that for certain graph topologies, the value of LC- k at node i is close to 1, implying that δ_i is an approximately k -smooth function on the graph. Although LC- k cannot be computed locally, methods exist for computing the value approximately, without having to perform an eigendecomposition of a potentially large graph Laplacian matrix [18], [27].

III. ALGORITHMS

We now describe the local algorithms that we will compare in this paper. A theoretical analysis of the relative convergence rates is provided in the next section.

A. Vanilla Random Walk

Algorithm 1 Vanilla RW

Require: Local access to graph G (i.e., at every node, all its neighbors are accessible) and number of iterations T .

- 1) **initialize:**
 - $i \leftarrow \text{Unif}\{1, \dots, n\}$ (Uniform sampling from V)
 - $i_{\max} \leftarrow i$
 - $f_{\max} \leftarrow f_i$
 - 2) **repeat** for T iterations:
 - $i \leftarrow \text{Unif}\{j : w_{ij} = 1\}$
 - if** $f_i > f_{\max}$:
 - $f_{\max} \leftarrow f_i$
 - $i_{\max} \leftarrow i$
 - return** $i_{\max}$
-

We begin by formally writing out the algorithm for an unbiased random walk on the vertices of the graph. The random walk proceeds around the nodes of the graph that are connected by edges. The transition kernel for the vanilla random walk is given by $P = D^{-1}W$. Thus, the transition probability from node i to node j is given by $P_{ij} = \frac{w_{ij}}{d_i}$,

where $w_{ij} = 1$ if an edge exists between i and j , and $w_{ij} = 0$ otherwise.

The vanilla random walk is agnostic to the function f , and converges to a stationary distribution that is proportional to the degrees of the graph nodes [28]. Note, however, that we cannot immediately conclude that such a function-agnostic random walk will not optimize the function efficiently—the class of k -smooth functions on graphs does not have a simple description, and it may happen that optimizers of such functions lie at high-degree nodes, which are precisely those nodes that attract the vanilla random walk.

B. Exponentially-Weighted Walk

The second algorithm involves biasing the choices of each move in the vanilla random walk according to the function values on the neighboring nodes. We choose an exponential weight function, so that the target probability density is defined by $p_f(i) \propto \exp(\gamma f_i)$. Starting with a proposal distribution $Q = D^{-1}W$, we use the MH method in equation (1) to construct the transition kernel for the exponentially-weighted walk, as follows:

$$\mathbf{P}_{ij} = \begin{cases} \frac{1}{d_i} w_{ij} R(i, j) & j \neq i, \\ 1 - \sum_{j \neq i} \mathbf{P}_{ij} & j = i, \end{cases} \quad (5)$$

where

$$R(i, j) = \min \left(1, \frac{p_f(j)(1/d_j)}{p_f(i)(1/d_i)} \right) = \min \left(1, e^{\gamma(f_j - f_i)} \frac{d_i}{d_j} \right).$$

Thus, the transition probability between any two nodes in the exponentially-weighted walk depends on the difference in function value $f_i - f_j$ and the degrees of the nodes. The parameter γ determines the “peakiness” of the target density p_f .

Unlike in the vanilla random walk, we assume that at every node, it is possible to access the function value and the degree of its neighbors. However, we do not make use of the smoothness constraint on the function. The parameter γ may be viewed as controlling the greedy nature of this algorithm. When $\gamma = +\infty$, the random walk always ascends; i.e., if $f(j) < f(i)$, then $P(i, j) = 0$. However, such a greedy walk is susceptible to getting stuck at local maxima and failing to find the global maxima quickly. At the other extreme, when $\gamma = 0$, the exponential random walk is function-agnostic (as in the vanilla random walk), and converges to the uniform distribution over vertices.

C. Graph Laplacian Walk

For the third algorithm, we use a proposal distribution that is derived from the graph Laplacian, instead of the vanilla random walk proposal distribution $Q = D^{-1}W$ that we used for the exponentially-weighted walk.

Let U_k represent the matrix containing k eigenvectors corresponding to the smallest k eigenvalues of the graph Laplacian. Let $\delta_i \in \{0, 1\}^n$ represent the indicator vector for node i . The proposal distribution for the graph Laplacian

Algorithm 2 Exponential RW(γ)

Require: Local access to graph G and function f (i.e., at every node, all its neighbors and their function values are accessible), number of iterations T , and tuning parameter γ .

- 1) **initialize:**
 - $i \leftarrow \text{Unif}\{1, \dots, n\}$ (Uniform sampling from V)
 - $i_{\max} \leftarrow i$
 - $f_{\max} \leftarrow f_i$
 - 2) **repeat** for T iterations:
 - Generate j according to $\mathbf{P}_{ij}$ given in equation (5)
 - $i \leftarrow j$
 - if** $f_i > f_{\max}$:
 - $f_{\max} \leftarrow f_i$
 - $i_{\max} \leftarrow i$
 - return** $i_{\max}$
-

walk is given by

$$\mathbf{Q}'_{ij} = \frac{\|U_k^T \delta_j\|_2^2}{\sum_{j: w_{ij}=1} \|U_k^T \delta_j\|_2^2} w_{ij}. \quad (6)$$

Note that the Laplacian-based proposal distribution Q' is indeed “local,” since $Q'_{ij} = 0$ whenever $Q_{ij} = 0$. For the Laplacian walk, we choose a target probability density according to $p_f(i) \propto f_i^2$. The squaring of the function restricts our analysis to maximizing positive functions. Using the MH method in equation (1), we derive the transition kernel for graph Laplacian walk, as follows:

$$\mathbf{P}_{ij} = \begin{cases} Q'_{ij} R(i, j) & j \neq i, \\ 1 - \sum_{j \neq i} \mathbf{P}_{ij} & j = i, \end{cases} \quad (7)$$

where we set $R(i, j) = 0$ if $w_{ij} = 0$; and if $w_{ij} = 1$, we have

$$\begin{aligned} R(i, j) &= \min \left(1, \frac{p_f(j) \mathbf{Q}'_{ji}}{p_f(i) \mathbf{Q}'_{ij}} \right) \\ &= \min \left(1, \frac{f_j^2 \|U_k^T \delta_i\|_2^2 \sum_{j: w_{ij}=1} \|U_k^T \delta_j\|_2^2}{f_i^2 \|U_k^T \delta_j\|_2^2 \sum_{i: w_{ji}=1} \|U_k^T \delta_i\|_2^2} \right). \end{aligned} \quad (8)$$

An important difference between the exponentially-weighted walk and the graph Laplacian walk is that the former algorithm is truly a local algorithm. Indeed, only the function values of the neighboring nodes are needed to determine the probability distribution for the next step of Algorithm 2. This differs from Algorithm 3, for which we assume knowledge of the graph Laplacian (or at least its top k eigenvectors) in order to compute $\|U_k^T \delta_i\|_2^2$ for the neighboring nodes of every iterate.

We now describe a variant of Algorithm 3 applicable to approximately smooth functions. We begin with a definition:

Definition 3. A function f is ϵ -approximately k -smooth if for the decomposition $f = f_{ks} + f_r$ where $f_{ks} = U_k^T U_k f$, we have

$$|f_r(i)| \leq \epsilon \|f_{ks}\|_2,$$

Algorithm 3 Laplacian RW(k)

Require: Function smoothness k , local access to graph G , function f and LC- k (i.e., at every node, all its neighbors, their function values, and their LC- k values are accessible/computable), and number of iterations T

```
1) initialize:
    $i \leftarrow \text{Unif}\{1, \dots, n\}$  (Uniform sampling from  $V$ )
    $i_{\max} \leftarrow i$ 
    $f_{\max} \leftarrow f_i$ 
2) repeat for  $T$  iterations:
   Generate  $j$  according to  $\mathbf{P}_{ij}$  given in equation (7)
    $i \leftarrow j$ 
   if  $f_i > f_{\max}$ :
      $f_{\max} \leftarrow f_i$ 
      $i_{\max} \leftarrow i$ 
return  $i_{\max}$ 
```

for all $i \in V$ and some $\epsilon > 0$.

For the class of ϵ -approximate k -smooth functions, the proposal distribution for the graph Laplacian walk given in equation (6) is replaced by

$$Q'_{ij} = \frac{(\|U_k^T \delta_j\|_2 + \epsilon)^2}{\sum_{j: w_{ij}=1} (\|U_k^T \delta_j\|_2 + \epsilon)^2} w_{ij}. \quad (9)$$

The transition kernel for the Laplacian walk will remain the same as in equation (7), with the modified Q'_{ij} in equation (9).

IV. THEORETICAL ANALYSIS

In this section, we provide results concerning the convergence rate of the local search algorithms. For both of our algorithms, we derive bounds on (i) the total variation distance between the probability distribution at time t and the stationary distribution of the Markov chain; and (ii) bounds on the hitting time of the algorithm in expectation and in high probability. A practical consequence of the total variation bounds is that if we run the corresponding local algorithms for sufficiently many steps, we are guaranteed that taking the maximizer of the empirically constructed distribution will be provably close to the function maximizer. The hitting time bounds may be interpreted as a bound on the amount of time needed to first visit a maximum—thus, if we halt the algorithms after a prescribed number of steps and choose the node with the largest function value so far, we are guaranteed to obtain a maximum with high probability. In fact, if we are given slightly more knowledge (i.e., that we have located a function maximizer upon visiting it for the first time), the theorems provide stronger guarantees for a variant of the local algorithms that halt once they identify a maximum.

A. Exponentially-Weighted Walk

In this section, we analyze the exponentially-weighted random walk in Algorithm 2.

1) *Convergence:* In the first theorem, we show the convergence of the exponentially-weighted random walk to the target density p_f in total variation norm.

Theorem 1 (Convergence of Algorithm 2). *For a connected graph G with diameter r , and for a graph function f that is nonzero on all vertices, the rate of convergence of the random walk proposed in Algorithm 3 to its stationary distribution p_f is given by*

$$\|\mathbf{P}_{i*}^t - p_f\|_{TV} \leq \left(1 - \frac{\delta_f^{r-1}}{(d_{\max} \Delta_f)^r}\right)^{\lfloor \frac{t}{r} \rfloor}, \quad \forall i \in V, \quad (10)$$

where $\mathbf{P}_{i*}^t$ is the i^{th} row of $\mathbf{P}^t$ (the transition probability matrix after t steps), $\delta_f = \min_i p_f(i)$, $\Delta_f = \max_i p_f(i)$.

Proof: We refer the reader to the full version of the paper: <https://arxiv.org/abs/1802.04475>.

2) *Results on Hitting Times:* In the next theorem, we prove an upper bound on the expected number of steps it takes for Algorithm 2 to reach the function maximum.

Theorem 2 (Expected hitting time). *For Algorithm 2, let $v_t \in V$ be the state at time t , and define $T_{\text{hit}} = \min\{t \geq 0 : f_{v_t} = f_{\max}\}$ to be the number steps it takes for Algorithm 2 to reach the function maximum. Let $f_{\min}$ denote the minimum value of the function f . Then the expected value of the hitting time is bounded by*

$$\mathbb{E}T_{\text{hit}} \leq d_{\max}^r e^{\gamma(r-1)(f_{\max}-f_{\min})}. \quad (11)$$

Proof: We refer the reader to the full version of the paper: <https://arxiv.org/abs/1802.04475>.

We now derive a high-probability bound on the hitting time.

Theorem 3 (High-probability bound on hitting time). *For any initial distribution μ and any $s > 0$, we have*

$$\mathbb{P}_{\mu}[T_{\text{hit}} > t] \leq \left(\frac{t_{\text{hit}}^*}{s}\right)^{\lfloor t/s \rfloor}, \quad (12)$$

where $t_{\text{hit}}^* = d_{\max}^r e^{\gamma(r-1)(f_{\max}-f_{\min})}$ and $t > 0$.

Proof: For any integer $m \geq 1$, we have

$$\begin{aligned} \mathbb{P}_{\mu}[T_{\text{hit}} > ms \mid T_{\text{hit}} > (m-1)s] &\leq \max_j \mathbb{P}_j(T_{\text{hit}} > s) \\ &\leq \frac{\mathbb{E}_j(T_{\text{hit}})}{s} = \frac{t_{\text{hit}}^*}{s}. \end{aligned}$$

By induction on m , we obtain $\mathbb{P}_{\mu}[T_{\text{hit}} > ms] \leq \left(\frac{t_{\text{hit}}^*}{s}\right)^m$, which implies inequality (12).

B. Graph Laplacian Walk

We begin by proving a few lemmas concerning the proposal distribution Q' used in Algorithm 3 for k -smooth and ϵ -approximately k -smooth functions. We establish the following envelope condition:

$$\frac{1}{M} p_f(j) \leq Q'_{ij}, \quad (13)$$

for an appropriate constant M , when $w_{ij} = 1$.

Lemma 1 (Dominance for k -smooth graph functions). *Suppose f is k -smooth. For the proposal distribution Q' given in equation (6), the envelope condition (13) holds with $M = k$.*

Proof: We refer the reader to the full version of the paper: <https://arxiv.org/abs/1802.04475>.

Lemma 2 (Dominance for ϵ -approximate k -smooth graph functions). *Suppose f is ϵ -approximate k -smooth. For the proposal distribution Q' given in equation (9), the envelope condition (13) holds with*

$$M = k + 2k\sqrt{n\epsilon} + n\epsilon^2. \quad (14)$$

Proof: We refer the reader to the full version of the paper: <https://arxiv.org/abs/1802.04475>.

1) *Convergence:* The first theorem concerns convergence in TV distance.

Theorem 4 (Convergence of Algorithm 3). *For a connected graph G with diameter r , and for a k -smooth graph function f (either exact or ϵ -approximate) that is nonzero on all the vertices, the rate of convergence of the random walk proposed in Algorithm 3 to its stationary distribution p_f is given by*

$$\|\mathbf{P}_{i*}^t - p_f\|_{TV} \leq \left(1 - \frac{\delta_f^{r-1}}{M^r}\right)^{\lfloor \frac{t}{r} \rfloor}, \quad \forall i \in V, \quad (15)$$

where $\mathbf{P}_{i*}^t$ is the i th row of $\mathbf{P}^t$, $\delta_f = \min_i p_f(i)$, and M is the dominance constant established in Lemmas 1 and 2.

Proof: We refer the reader to the full version of the paper: <https://arxiv.org/abs/1802.04475>.

2) *Results on Hitting Times:* In the next theorem, we prove an upper bound on the expected number of steps it takes for Algorithm 3 to reach the function maximum.

Theorem 5 (Expected hitting time). *For Algorithm 3, suppose $v_t \in V$ be the state at time t , and define $T_{hit} = \min\{t \geq 0 : f_{v_t} = f_{max}\}$. The expected value of the hitting time is bounded by*

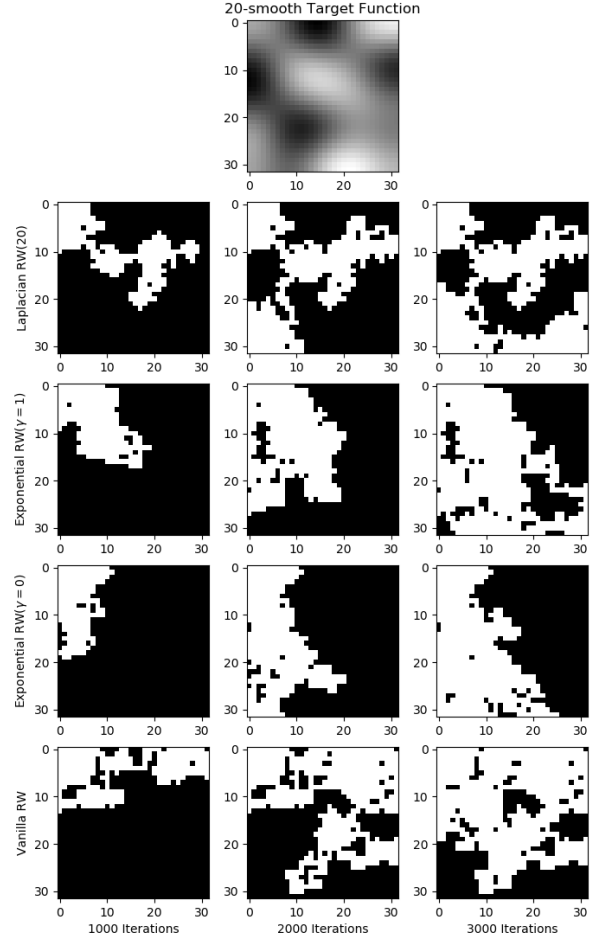
$$\mathbb{E}T_{hit} \leq \frac{(M\|f\|^2)^r}{f_{max}^2 f_{min}^{2(r-1)}}. \quad (16)$$

Proof: We refer the reader to the full version of the paper: <https://arxiv.org/abs/1802.04475>.

The high-probability bound on expected time would be same as that in Theorem 3, with

$$t_{hit}^* = \frac{(M\|f\|^2)^r}{f_{max}^2 f_{min}^{2(r-1)}}.$$

Fig. 2. An example run of the algorithms on a 20-smooth function on a 32×32 2D grid graph. White pixels in rows 2 to 5 indicate nodes that have been visited by the algorithm. White pixels in the function plot denote high function values. Notice that the Laplacian RW is very effective at traversing regions of the graph where the function takes large values. The exponential RW with $\gamma = 1$ is also effective at reaching the function peaks located at the coordinates (15,10). The vanilla random walk covers a large area but without any preferred direction, as expected.



V. EXPERIMENTS

In this section, we consider various graph topologies and approximately k -smooth functions defined on the nodes of the graphs and compare the algorithms described in Section III. We simulated Algorithms 1, 2 and 3 on the following graph models: (i) 2D Grid graph, (ii) Erdős-Renyi (ER) graph, (ii) and Barabasi-Albert (BA) graph. For the 2D grid graph, we ran the algorithms on a 32×32 grid of $n = 1024$ nodes. For the ER graph, we generated a random graph with $n = 1000$ nodes with the probability of an edge between any two nodes being $p = \frac{1.1 \log n}{n}$, and discarded the graphs with isolated nodes so that the entire graph formed one connected component. For the BA graph, we generated a random graph with $n = 1000$ nodes starting with a seed set of $m = 3$ nodes.

For each of these graphs, we generated random smooth functions for varying values of k . For each value of k , we sampled α from a $k \times 1$ vector of standard normal random variables, and constructed the smooth function as

$f = U_k \alpha$, where $U \in \mathbb{R}^{n \times k}$ is the matrix containing the top k eigenvectors of the graph Laplacian matrix. The functions were then made nonnegative by lifting the function value at every node by the minimum. Since the constant vector is the first eigenvector of the graph Laplacian, this process of adding a constant vector does not affect the smoothness of the function.

The algorithms were run till the random walk “hit” the maximum, or up to a maximum of 10,000 steps. The hitting times were then averaged over 100 iterations for each value of k , and for 10 randomly chosen functions (i.e., a total of 1000 iterations for each value of k). For Algorithm 2, we ran the experiments with two different values of γ . The case $\gamma = 0$ corresponds to a function-agnostic random walk that converges to a uniform stationary distribution across all the nodes. The case $\gamma = 1$ corresponds to a moderately greedy random walk that converges to a stationary distribution proportional to $e^{\gamma f}$. A high value of γ would reduce Algorithm 2 to a ‘greedy’ walk that gets stuck in local minima, which is not desirable. For instance, in our experiments, $\gamma = 10$ led to very poor performance due the algorithm often being unable to move out of local minima.

Figure 3 reports the results for various random graph models. As can be seen, the Laplacian walk of Algorithm 3 consistently outperformed all other algorithms for all the three classes of graphs we considered. It is also worth noting that changing the comparison criterion from hitting the global maximum to hitting the top 1% of nodes did not change the performance of the algorithms in our simulations.

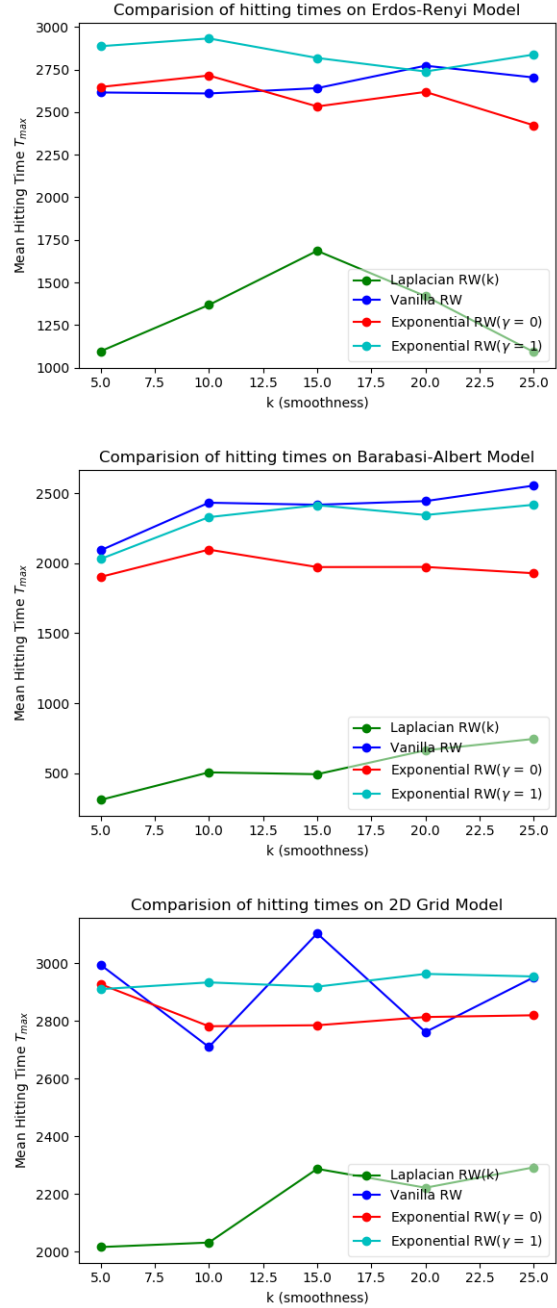
VI. DISCUSSION

We have presented two new algorithms for graph function maximization based on local movements around the nodes of a graph. Our first algorithm concerns an exponentially weighted random graph, and our second algorithm uses information about the spectrum of the graph Laplacian matrix. We have provided theoretical results concerning the rate of convergence of our algorithms in terms of total variation distance and hitting time when the graph function belongs to a certain smoothness class.

The algorithms we have studied in this paper only involve local movements along edges, from one node to an adjoining neighbor. However, one might imagine variants of the algorithms that allow “jump” movements that reinitialize the walk at a randomly drawn vertex, either uniformly selected from all the nodes or all the nodes previously explored in previous time steps. It would be interesting to see how incorporating the option of a jump move might affect the convergence analysis of the algorithms.

Another question to explore is the tightness of our bounds on rates of convergence of the algorithms. Our results are derived based on a seemingly coarse analysis, and it would be interesting to see if it is possible to find worst-case graphs and classes of smooth functions for which one can also derive lower bounds for the rate of convergence of any local algorithm. This also gives rise to the important related

Fig. 3. Comparison of hitting times as a function of smoothness for various random graph models



question of the landscape of local and global optima of a k -smooth function, where the smoothness is defined in terms of eigenvalues of the graph Laplacian. This appears to be a challenging problem even for $k = 2$ and for Erdős-Renyi random graphs. Ideally, we would also like to be able to translate the bounds on total variation distance and hitting time into precise recommendations regarding the number of iterates required for each of our algorithms to locate a maximum.

A final question concerns improving the convergence rate of local algorithms when the next iterate is allowed to depend

on the values of the function in a neighborhood of radius r around the current iterate. The algorithms described in this paper are limited to the case when $r = 1$. However, if the local algorithm were given information about a larger neighborhood at each step, perhaps it would be possible to devise an analog of higher-order descent algorithms, which are known to exhibit faster rates of convergence in continuous optimization settings.

REFERENCES

- [1] D. Chen, L. Lü, M.-S. Shang, Y.-C. Zhang, and T. Zhou, “Identifying influential nodes in complex networks,” *Physica A: Statistical Mechanics and Its Applications*, vol. 391, no. 4, pp. 1777–1787, 2012.
- [2] M. O. Jackson, *Social and Economic Networks*. Princeton University Press, 2010.
- [3] N. K. Vishnoi, “ $lx = b$ laplacian solvers and their algorithmic applications,” *Foundations and Trends in Theoretical Computer Science*, vol. 8, no. 1–2, pp. 1–141, 2013.
- [4] R. I. Kondor and J. Lafferty, “Diffusion kernels on graphs and other discrete input spaces,” in *ICML*, vol. 2, 2002, pp. 315–322.
- [5] A. J. Smola and R. Kondor, “Kernels and regularization on graphs,” in *Learning Theory and Kernel Machines*. Springer, 2003, pp. 144–158.
- [6] Y.-X. Wang, J. Sharpnack, A. J. Smola, and R. J. Tibshirani, “Trend filtering on graphs,” *Journal of Machine Learning Research*, vol. 17, no. 105, pp. 1–41, 2016.
- [7] M. Brautbar and M. J. Kearns, “Local algorithms for finding interesting individuals in large networks,” 2010.
- [8] D. D. Heckathorn, “Respondent-driven sampling: A new approach to the study of hidden populations,” *Social Problems*, vol. 44, no. 2, pp. 174–199, 1997.
- [9] A. Banerjee, A. G. Chandrasekhar, E. Duflo, and M. O. Jackson, “Gossip: Identifying central individuals in a social network,” National Bureau of Economic Research, Tech. Rep., 2014.
- [10] D. P. Bertsekas, *Nonlinear Programming*. Athena Scientific Belmont, 1999.
- [11] M. Welling and Y. W. Teh, “Bayesian learning via stochastic gradient Langevin dynamics,” in *Proceedings of the 28th International Conference on Machine Learning*, 2011, pp. 681–688.
- [12] Y. Nesterov, “A method of solving a convex programming problem with convergence rate $o(1/k^2)$,” in *Soviet Mathematics Doklady*, vol. 27, no. 2, 1983, pp. 372–376.
- [13] —, *Introductory Lectures on Convex Optimization: A Basic Course*. Springer Science & Business Media, 2013, vol. 87.
- [14] Y. Lin, L. Lu, and S.-T. Yau, “Ricci curvature of graphs,” *Tohoku Mathematical Journal, Second Series*, vol. 63, no. 4, pp. 605–627, 2011.
- [15] H. Hirai, “Discrete convex functions on graphs and their algorithmic applications,” in *Combinatorial Optimization and Graph Algorithms*. Springer, 2017, pp. 67–100.
- [16] D. I. Shuman, S. K. Narang, P. Frossard, A. Ortega, and P. Vandergheynst, “The emerging field of signal processing on graphs: Extending high-dimensional data analysis to networks and other irregular domains,” *IEEE Signal Processing Magazine*, vol. 30, no. 3, pp. 83–98, 2013.
- [17] A. Anis, A. Gadde, and A. Ortega, “Efficient sampling set selection for bandlimited graph signals using graph spectral proxies,” *IEEE Transactions on Signal Processing*, vol. 64, no. 14, pp. 3775–3789, 2016.
- [18] G. Puy, N. Tremblay, R. Gribonval, and P. Vandergheynst, “Random sampling of bandlimited signals on graphs,” *Applied and Computational Harmonic Analysis*, 2016.
- [19] P. R. Christian and G. Casella, *Monte Carlo Statistical Methods*. Springer New York, 1999.
- [20] A. S. Dalalyan, “Theoretical guarantees for approximate sampling from smooth and log-concave densities,” *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, vol. 79, no. 3, pp. 651–676, 2017.
- [21] Y. Zhang, P. Liang, and M. Charikar, “A hitting time analysis of stochastic gradient Langevin dynamics,” *arXiv preprint arXiv:1702.05575*, 2017.
- [22] M. Raginsky, A. Rakhlin, and M. Telgarsky, “Non-convex learning via stochastic gradient Langevin dynamics: a nonasymptotic analysis,” *arXiv preprint arXiv:1702.03849*, 2017.
- [23] C. Borgs, M. Brautbar, J. Chayes, S. Khanna, and B. Lucier, “The power of local information in social networks,” in *International Workshop on Internet and Network Economics*. Springer, 2012, pp. 406–419.
- [24] A. Frieze and W. Pegden, “Looking for vertex number one,” *The Annals of Applied Probability*, vol. 27, no. 1, pp. 582–630, 2017.
- [25] G. Grimmett and D. Stirzaker, *Probability and Random Processes*, ser. Probability and Random Processes. OUP Oxford, 2001.
- [26] U. Von Luxburg, “A tutorial on spectral clustering,” *Statistics and Computing*, vol. 17, no. 4, pp. 395–416, 2007.
- [27] N. Tremblay, G. Puy, R. Gribonval, and P. Vandergheynst, “Compressive spectral clustering,” in *International Conference on Machine Learning*, 2016, pp. 1002–1011.
- [28] D. Aldous and J. Fill, “Reversible Markov chains and random walks on graphs,” 2002.