

A continuum limit for the PageRank algorithm

A. YUAN¹, J. CALDER¹ and B. OSTING²

¹*Department of Mathematics, University of Minnesota, Minneapolis, MN 55455, USA*
emails: yuanx290@umn.edu; jwcalder@umn.edu

²*Department of Mathematics, University of Utah, Salt Lake City, UT 84112, USA*
email: osting@math.utah.edu

(Received 9 April 2020; revised 10 January 2021; accepted 26 March 2021)

Semi-supervised and unsupervised machine learning methods often rely on graphs to model data, prompting research on how theoretical properties of operators on graphs are leveraged in learning problems. While most of the existing literature focuses on undirected graphs, directed graphs are very important in practice, giving models for physical, biological or transportation networks, among many other applications. In this paper, we propose a new framework for rigorously studying continuum limits of learning algorithms on directed graphs. We use the new framework to study the PageRank algorithm and show how it can be interpreted as a numerical scheme on a directed graph involving a type of normalised *graph Laplacian*. We show that the corresponding continuum limit problem, which is taken as the number of webpages grows to infinity, is a second-order, possibly degenerate, elliptic equation that contains reaction, diffusion and advection terms. We prove that the numerical scheme is consistent and stable and compute explicit rates of convergence of the discrete solution to the solution of the continuum limit partial differential equation. We give applications to proving stability and asymptotic regularity of the PageRank vector. Finally, we illustrate our results with numerical experiments and explore an application to data depth.

Key words: Partial differential equations on graphs and networks, second-order elliptic equations, viscosity solutions

2020 Mathematics Subject Classification: 35J15 (Primary); 35D40 (Secondary)

1 Introduction

Due to its versatility in modelling data, graphs are frequently leveraged for applications in machine learning and data science. A graph structure encodes interdependencies among constituents, such as social media users, images or videos in a collection, or physical or biological agents, and provides a convenient representation for high dimensional data. For example, in a graph representing research collaborations, we can represent each author as a node in the graph, and co-authorship is represented by edges between nodes, with edge weights depending on the frequency of co-authorship. The resulting graph is *undirected* as the edges are bi-directional. On the other hand, transportation and biological networks often result in *directed* graphs because the relationship between two nodes is ordered, such as the direction of a train route or a predator–prey relationship.

Graph-based methods are particularly prominent in unsupervised and semi-supervised machine learning tasks that seek to reveal structures and patterns in unlabelled data. For example, in semi-supervised classification, one has labels for a subset of the nodes in the graph, and the problem is to propagate the labels to the rest of the graph in a meaningful way. A widely used and very successful algorithm for semi-supervised classification is Laplacian semi-supervised learning, originally proposed in [61], which finds the unique *graph harmonic* function that extends the labels. There are many extensions and modifications of Laplacian regularisation (see, e.g., [58, 56, 59, 1, 48, 3]), with more recent methods drawing inspiration from partial differential equations (PDEs) [6, 23]. For classification problems at very low labelling rates, p -Laplacian regularisation has recently been introduced [20, 22]. In unsupervised learning, graph-based algorithms are used in spectral clustering [40, 42], Laplacian eigenmaps [2], diffusion maps [17], manifold ranking [32, 33, 52, 55, 60, 54], minimal surface graph partitioning [62] and PageRank [30].

Various types of graph Laplacians appear in nearly all graph-based learning algorithms, due to the ability of the graph Laplacian to uncover geometric structure in datasets. Graph Laplacians that are commonly used in practice include the *unnormalised* Laplacian

$$\mathcal{L}u(x) = \sum_{y \in X} \omega_{xy}(u(y) - u(x)),$$

the *random walk* Laplacian

$$\mathcal{L}^{rw}u(x) = \frac{1}{d_x} \sum_{y \in X} \omega_{xy}(u(y) - u(x))$$

and the *normalised* Laplacian

$$\mathcal{L}^nu(x) = \sum_{y \in X} \frac{\omega_{xy}}{\sqrt{d_x d_y}} u(y) - u(x),$$

where X denotes the set of nodes in the graph, $u: X \rightarrow \mathbb{R}$, ω_{xy} is the (undirected) edge weight between x and y , and $d_x = \sum_{y \in X} \omega_{xy}$ is the degree of node x . The unnormalised graph Laplacian appears naturally as the gradient of the Dirichlet energy

$$E(u) = \sum_{x,y \in X} \omega_{xy}(u(x) - u(y))^2.$$

The random walk Laplacian is exactly the generator for a random walk on X with probability $d_x^{-1}\omega_{xy}$ of stepping from x to y , and the normalised graph Laplacian is a convenient way to obtain a *symmetric* normalisation of the graph Laplacian. While these normalisations are most frequently used in practice, many other choices are possible. For example, see [36] for an analysis of how the choice of normalisation affects spectral clustering. We note that the random walk interpretation allows us to view methods like those studied in [61] as performing classification by randomly walking on the graph until hitting a labelled node. Intuitively, the random walk will naturally *learn* the structure of the unlabelled data by remaining within clusters of high density for long enough to hit a labelled point, before moving to a different cluster. While many classification algorithms seek graph harmonic functions, the spectrum of graph Laplacians is widely used to construct low dimensional embeddings of graphs.

The algorithms discussed above are mainly designed for symmetric graphs, where $\omega_{xy} = \omega_{yx}$. Perhaps one of the most widely known algorithms for directed graphs is the PageRank algorithm,

which is used to evaluate the importance of nodes in a graph based on their link structure. While the algorithm is most famous for sorting Google search results up until the mid-2000s, variants of PageRank are used by other tech companies (e.g., Twitter uses a reversed PageRank to identify influential, topic-specific accounts) and have been adapted to solve problems in neuroscience, genetics and recommender systems [30]. The PageRank algorithm uses a random surfer model with teleportation probability $\alpha \in [0, 1]$ to rank pages. To describe the model, when the random surfer is at webpage x , she will with probability α teleport to a random webpage, and with probability $1 - \alpha$ click on an outgoing link to another webpage. When the surfer clicks on an outgoing link, the link is selected at random and we denote by p_{xy} the probability of clicking a link to website y from website x . When she randomly teleports, the next website is chosen at random from a teleportation probability distribution $\mathbf{v}$. The inclusion of the teleportation step ensures the random surfer does not get stuck in disconnected components of the graph.

The PageRank vector is the invariant distribution of the resulting Markov chain, which measures the amount of time the random surfer spends on each webpage. Webpages that are visited more often by the surfer are ranked more highly, while websites that are rarely visited are ranked lower. Mathematically, the PageRank vector $\mathbf{r}$ is the (normalised) solution of the eigenvector problem

$$((1 - \alpha)P + \alpha \mathbf{v} \mathbf{1}^T) \mathbf{r} = \mathbf{r}, \quad (1.1)$$

where $P = (p_{yx})_{x,y \in \mathcal{X}}$ is the probability transition matrix described above, $\mathbf{v}$ is the teleportation probability distribution and $\mathbf{1}$ is the column vector of all 1's. We note that by the Perron–Frobenius theorem [30], the PageRank vector $\mathbf{r}$ can be chosen to have real-valued strictly positive entries. If we choose the normalisation $\mathbf{1}^T \mathbf{r} = 1$, so that $\mathbf{r}$ is a probability distribution, then the eigenvector problem (1.1) is equivalent to the linear system

$$(I - (1 - \alpha)P) \mathbf{r} = \alpha \mathbf{v}. \quad (1.2)$$

This formulation is more convenient, since the left-hand side can be interpreted as a type of graph Laplacian.

The teleportation probability distribution $\mathbf{v}$ can be uniform over all webpages or can be nonuniform. Indeed, by setting $\mathbf{v}(x) = \delta_{x_0}(x)$ for a specific website x_0 leads to a localised PageRank algorithm that ranks sites nearby x_0 [30]. Computationally, the PageRank vector is obtained via the power method on (1.1), which converges at a rate of $|\lambda_2/\lambda_1|$, that is, a ratio of the second eigenvalue to the leading eigenvalue of the matrix. In the case of PageRank, Haveliwala and Kamvar [31] show that $\lambda_1 = 1$ and $\lambda_2 = 1 - \alpha$, so the convergence rate depends heavily on the choice of the teleportation parameter. Google takes $1 - \alpha$ to be 0.85 [38]. There are also adaptations of semi-supervised learning to directed graphs (see [57]).

Due to the ubiquity of graph Laplacians in graph-based learning problems, much work has been devoted to understanding how the graph Laplacian is able to uncover geometric and distributional structure from unlabelled data. To do this, one usually assumes the graph is a random geometric graph with n points and length scale $h > 0$ and considers the limit as $n \rightarrow \infty$ and $h \rightarrow 0$. This means the nodes in the graph are an *i.i.d.* sample of size n from a density ρ supported on a d -dimensional manifold $\mathcal{M}$ embedded in $\mathbb{R}^D$, and the weights ω_{xy} are defined by

$$\omega_{xy} = \Phi \left(\frac{|x - y|}{h} \right),$$

where $\Phi : [0, \infty) \rightarrow [0, \infty)$ is nonincreasing and usually compactly supported. The first results to appear in the literature were pointwise consistency results, showing that a graph Laplacian $\mathcal{L}$ applied to a smooth test function $\varphi \in C^3(\mathcal{M})$ converges, as $n \rightarrow \infty$ and $h \rightarrow 0$ to a weighted version of the Laplace–Beltrami operator

$$\Delta_\rho \varphi = \rho^a \operatorname{div}(\rho^b \nabla(\rho^c \varphi))$$

for various values of a, b, c that depend on the choice of normalisation of the graph Laplacian. For example, for the unnormalised graph Laplacian, $a = -1, b = 2, c = 0$, and for the random walk Laplacian $a = -2, b = 2, c = 0$. If $h \rightarrow 0$ and $n \rightarrow \infty$ simultaneously, then the condition $nh^{d+2} \gg \log n$ is required for pointwise consistency, which ensures there are enough neighbours of each data point to apply appropriate concentration of measure results. To obtain $O(h)$ pointwise consistency rates, it is required that $nh^{d+4} \gg 1$. We contrast this with the condition $nh^d \gg \log n$ required for graph connectivity. For pointwise consistency results of this flavour, see [7, 37, 35, 34, 4, 46]. Pointwise consistency was extended to k -nearest neighbour graph constructions in [49], which includes some mildly directed graphs due to antisymmetries in the k -nearest neighbour relation.

While pointwise consistency results are informative, they do not prove that the solutions of graph-based problems converge to solutions of their counterparts as $n \rightarrow \infty$ and $h \rightarrow 0$. This question is more subtle and requires further analysis. The problem of spectral convergence of the graph Laplacian spectrum to that of the Laplace–Beltrami operator has been well-studied. Belkin and Niyogi [5] established L^2 spectral convergence (convergence of eigenvalues and L^2 convergence of eigenvectors) when ρ is the uniform distribution and this was extended to nonuniform distributions with partial convergence rates in [51]. Shi [43] proved convergence rates and extended the analysis to include manifolds with boundary. The L^2 convergence rate was improved recently in [25] using variational methods, and then further improved in [13] to agree with the pointwise consistency rate $O(h)$, which is the sharpest known L^2 spectral convergence rate. The variational parts of the spectral convergence arguments in [13, 25] were heavily influenced by earlier work in a non-probabilistic setting [8]. We also mention that very recent work has established the first L^∞ eigenvector convergence rates [19].

For problems in clustering and semi-supervised learning, recent work has begun to address convergence in the continuum using tools from the calculus of variations and viscosity solutions of PDEs. Trillos and Slepčev [50] developed a Gamma-convergence framework for proving discrete to continuum convergence of graph-based problems, and the framework has been applied to prove discrete to continuum consistency in many problems (see, e.g., [50, 28, 27, 24, 47, 41] and the references therein). Discrete to continuum convergence can also be established with the maximum principle and the viscosity solution framework, as was established in [9, 11] for the game-theoretic graph p -Laplacian and Lipschitz learning. The maximum principle can also be used to prove asymptotic Hölder regularity of solutions to graph-based learning problems, as was done in [14, 9]. For the linear 2-graph Laplacian, [26] used the maximum principle to establish discrete to continuum convergence rates for regression problems, and [15] used the maximum principle in combination with random walk arguments to establish convergence rates for semi-supervised learning at low labelling rates. We also mention that [44] uses the maximum principle to prove convergence rates for a reweighted version of the graph Laplacian in low label rate semi-supervised learning context.

Despite the flurry of recent work on discrete to continuum consistency results, almost none of the results apply to problems on *directed graphs*, which are important and widely used in practice. The only results we are aware of for directed graphs are for k -nearest neighbour graphs [49, 24], which are directed only due to the asymmetry of the k -nearest neighbour relation. Discrete to continuum results are important for providing insights and further understanding of algorithms. Furthermore, continuum limits allow us to prove stability of graph-based algorithms, showing that they are insensitive to the particular realisation of the data, and often can lead to new formulations of learning problems founded on stronger theoretical principles. This paper aims to start filling this void by studying consistency results for problems on *directed graphs*. We propose the *random directed geometric graph* model, which extends the random geometric graph in a natural way by adding directionality in the weights. For concreteness, we study the PageRank problem and prove that the PageRank vector converges in the large sample size limit to the solution of a continuum, possibly degenerate, elliptic PDE. Depending on the strength of the directionality in the weights, the continuum PDE can be a first-order equation, which is a new type of result for consistency of graph Laplacians. Our main results are finite sample size error estimates with high probability, which imply convergence in the continuum, but are stronger in that they hold in the non-asymptotic regime. We use these results to prove stability of the PageRank problem, and we also study the time-dependent version of the problem, which examines the continuum limit of the distribution of the random surfer. Our proofs use pointwise consistency and the maximum principle, with appropriate adaptations to directed graphs. We also present the results of some numerical experiments confirming our theoretical results and exploring applications to data depth.

2 Setup and main results

We now describe our setup and main results. Section 2.1 introduces our *random directed geometric graph* model, and Section 2.2 formulates the PageRank problem in a new way and gives our main results.

2.1 Random directed geometric graph

In order to study continuum limits for problems on directed graphs, we propose a new model for a random directed graph that we call a *random directed geometric graph*. Let $x_1, x_2, \dots, x_n$ be an *i.i.d.* sample of size n on the torus $\mathbb{T}^d = \mathbb{R}^d / \mathbb{Z}^d$ with a density function $\rho : \mathbb{T}^d \rightarrow [0, \infty)$. We define the weight ω_{xy} from x to y by

$$\omega_{xy} = \Phi \left(\frac{|B(x)(y - x - \varepsilon b(x))|}{h} \right),$$

where $B : \mathbb{T}^d \rightarrow \mathbb{R}^{d \times d}$ and $b : \mathbb{T}^d \rightarrow \mathbb{R}^d$, and $B(x)$ has full rank for every $x \in \mathbb{T}^d$. The parameter $h > 0$ is the bandwidth of the kernel, and $\varepsilon > 0$ is the strength of the directionality. The kernel function Φ is assumed to be nonnegative with compact support. When $B = I$ and $b = 0$ or $\varepsilon = 0$, the weights are the same as those of a random geometric graph, which is symmetric. For other choices of B and b , the corresponding graph is directed, with directional influence along the vector field b . The matrix B can be viewed as changing the metric locally.

Assume for the moment that Φ has compact support in $[0, 2]$. We observe that for fixed x , the support of the weight $y \mapsto \omega_{yx}$ is the ellipse shaped region

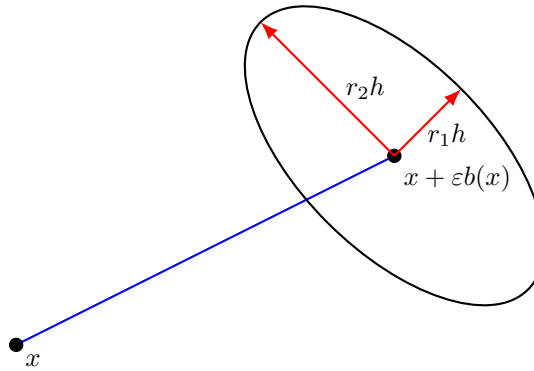


FIGURE 1. A visualisation of the set E_x , which is the support of the weight $y \mapsto \omega_{xy}$ in $\mathbb{R}^2$. The blue line represents the directional preference $b(y)$ multiplied by ε . The ellipse is the set E_x , where $r_i = 2/\sqrt{\lambda_i}$ and λ_1, λ_2 are the eigenvalues of $B(x)^T B(x)$, and the red arrows indicate the eigenvectors of $B(x)^T B(x)$.

$$E_x := \left\{ y \in \mathbb{R}^d \mid \frac{|B(x)(y - x - \varepsilon b(x))|}{h} < 2 \right\} \quad (2.1)$$

which depends on x . Figure 1 gives us a sense of E_x in two dimensions. In the random walk (or random surfer) interpretation, the random walker moves from x to a point in the set E_x , which contains a drift term $\varepsilon b(x)$ and an anisotropic diffusion governed by B .

In the following remarks, we provide an in-depth motivation for the weights in the context of ranking players in a sports tournament and modelling systematic distortion in the data [45].

Remark 2.1 (Motivation via ranking) *Assume each team or player is represented by a feature vector $x \in \mathbb{R}^d$ that adequately describes each player. When player x and player y play against each other, we write $x \succ y$ if x wins the game, and $y \succ x$ if y wins. We assume each time x and y play, x wins with probability $\mathbb{P}(x \succ y)$ and y wins with probability $\mathbb{P}(y \succ x) = 1 - \mathbb{P}(x \succ y)$.*

Suppose that x and y play n games, and we assign an edge in our graph from x to y if y wins more than half of the games, and assign the edge from y to x if x wins more than half the games. That is, the edge is directed towards the ‘better’ player. The weight on the edge is the excess number of wins for the winning player. In expectation, if $\mathbb{P}(x \succ y) \geq \frac{1}{2}$, then we have an edge from y to x with weight

$$\omega_{yx} = -1 + 2\mathbb{P}(x \succ y).$$

If $\mathbb{P}(y \succ x) \geq \frac{1}{2}$, then we have an edge from x to y with weight

$$\omega_{xy} = -1 + 2\mathbb{P}(y \succ x) = 1 - 2\mathbb{P}(x \succ y).$$

Now, we make a modelling assumption on $\mathbb{P}(x \succ y)$. We assume there is some underlying (unknown) ranking function

$$\varphi : \mathbb{R}^d \rightarrow [0, 1],$$

so that $\varphi(x) \geq \varphi(y)$ indicates player x is better than y . A natural model for the probability $\mathbb{P}(x \succ y)$ is then

$$\mathbb{P}(x \succ y) = \frac{1}{2} + \frac{\varphi(x) - \varphi(y)}{2},$$

leading to the weight

$$\omega_{yx} = \varphi(x) - \varphi(y),$$

when $\varphi(x) \geq \varphi(y)$. Using a Taylor expansion, we have an edge from y to x if

$$\varphi(y) \leq \varphi(x) \approx \varphi(y) + \nabla \varphi(y) \cdot (x - y) + \frac{1}{2}(x - y)^T \nabla^2 \varphi(x - y),$$

or

$$\nabla \varphi(y) \cdot (x - y) + \frac{1}{2}(x - y)^T \nabla^2 \varphi(x - y) \geq 0. \quad (2.2)$$

If φ is convex, the set of x satisfying the inequality above lie in an ellipse. This gives some motivation for the directional preference $b = \nabla \varphi$ and for the elliptical shape governed by $B = (\nabla^2 \varphi)^{1/2}$ as they occur in the definition of our weights. We restrict the weights locally to some ball $B(y, 2h)$ based on the assumption that teams play against similarly ranked teams in a tournament.

Remark 2.2 The work in [45] constructs two operators that can identify the common structures and the differences, respectively, between two diffeomorphic Riemannian manifolds. An application to identifying signals from foetal electrocardiogram data via observed maternal ECG data is considered. Their algorithm handles cases where there is a systematic diffusion in the observed vs. target data; our problem is ‘adjacent’ in the sense that we model the diffusion and directional preferences via $B(y)$ and $b(y)$ in our setup.

2.2 Main results

We now present our setup and main results. We take the random directed geometric graph model from Section 2.1 with $B(x) \equiv I$ (though see Section for a discussion of how the results change when B is not the identity). That is, let $x_1, x_2, \dots, x_n$ be an *i.i.d.* sample of size n on the torus $\mathbb{T}^d$ with probability density $\rho : \mathbb{T}^d \rightarrow [\rho_{\min}, \infty)$, where $\rho_{\min} > 0$, and set $X_n = \{x_1, x_2, \dots, x_n\}$. We define the weight from x to y by

$$\omega_n(x, y) = \Phi \left(\frac{|y - x - \varepsilon b(x)|}{h} \right), \quad (2.3)$$

and the degree of x by $d_n(x) = \sum_{y \in X_n} \omega_n(x, y)$. We assume the kernel Φ is smooth, compactly supported on $[0, 2]$, nonnegative, nonincreasing, satisfies $\Phi(0) > 0$ and

$$\int_{B(0,2)} \Phi(|z|) dz = 1. \quad (2.4)$$

As constructed, for example in [48], the probability of a random walk on the graph transitioning from x to y is $p_{xy} = d_n(x)^{-1} \omega_n(x, y)$. Plugging this into (1.2) and denoting the PageRank vector by $r_n : X_n \rightarrow \mathbb{R}$, we find that r_n satisfies the linear system

$$r_n(x) - (1 - \alpha) \sum_{y \in X_n} \frac{\omega_n(y, x)}{d_n(y)} r_n(y) = \alpha v(x) \quad \text{for all } x \in X_n, \quad (2.5)$$

where $\alpha \in (0, 1]$ is the teleportation probability and $v(x)$ is the teleportation probability distribution. To simplify the problem, we consider the normalised PageRank vector

$$u_n(x) = \frac{nh^d}{d_n(x)} r_n(x). \quad (2.6)$$

The degree $d_n(x)$ is the most basic measure of the importance of a node in a graph, and the normalised PageRank vector factors out the direct dependence on the degree to give an understanding of the additional geometric structure uncovered by PageRank. We easily see that the normalised PageRank vector $u_n : X_n \rightarrow \mathbb{R}$ satisfies the equation

$$d_n(x)u_n(x) - (1 - \alpha) \sum_{y \in X_n} \omega_n(y, x)u_n(y) = \alpha nh^d v(x) \quad \text{for all } x \in X_n. \quad (2.7)$$

We note that (2.7) is considerably simpler to analyse than (2.5) since the degree term $d_n(y)$ does not appear inside the summation. We rewrite this equation by defining the PageRank operator

$$\mathcal{L}_n u(x) := \frac{1}{d_n(x)} \sum_{y \in X_n} \omega_n(y, x)u(y) - u(x). \quad (2.8)$$

Then (2.7) can be written as

$$u_n(x) - \gamma \mathcal{L}_n u_n(x) = \frac{nh^d}{d_n(x)} v(x) \quad \text{for all } x \in X_n, \quad (2.9)$$

where $\gamma = (1 - \alpha)/\alpha$. We note that when the graph is symmetric, the PageRank operator is exactly the random walk graph Laplacian.

The corresponding problem in the continuum is the, possibly degenerate, elliptic PDE

$$u + \gamma_\varepsilon \rho^{-2} \operatorname{div}(\rho^2 b u) - \frac{1}{2} \sigma_\Phi \gamma_h \rho^{-2} \operatorname{div}(\rho^2 \nabla u) = \rho^{-1} v \quad \text{on } \mathbb{T}^d, \quad (2.10)$$

where $\sigma_\Phi = \int_{\mathbb{R}^d} \Phi(|z|) z_1^2 dz$,

$$\gamma_\varepsilon = \frac{(1 - \alpha)\varepsilon}{\alpha} \quad \text{and} \quad \gamma_h = \frac{(1 - \alpha)h^2}{\alpha}. \quad (2.11)$$

We also denote

$$\eta = \|\rho^{-2} \operatorname{div}(\rho^2 b)\|_{L^\infty(\mathbb{T}^d)}. \quad (2.12)$$

Our first main result is the following continuum limit.

Theorem 2.3 (Convergence of the Second-Order PageRank Problem) *Let $\rho \in C^{2,v}(\mathbb{T}^d)$, $b \in C^{2,v}(\mathbb{T}^d; \mathbb{R}^d)$ and $v \in C^{1,v}(\mathbb{T}^d)$ for some $0 < v < 1$. Assume that $\gamma_\varepsilon \leq 1$, $0 < \gamma_h \leq 1$, and $\eta < 1$. Let u_n be the solution to the PageRank problem (2.9) and let $u \in C^3(\mathbb{T}^d)$ be the solution to the*

PDE (2.10). Then there exists $C_1, C_2, c_1, c_2 > 0$ with C_1 depending on $\gamma_h > 0$, such that when $\varepsilon + h \leq c_1(1 - \eta\gamma_\varepsilon)$ we have that

$$\max_{x \in \mathcal{X}_n} |u(x) - u_n(x)| \leq C_1(1 - \eta\gamma_\varepsilon)^{-1}(\lambda + \varepsilon + h) \quad (2.13)$$

holds with probability at least $1 - C_2n \exp(-c_2nh^{d+2}\lambda^2) - C_2n \exp(-c_2nh^{d+2}(1 - \eta\gamma_\varepsilon)^2)$, where $0 < \lambda \leq 1$.

Remark 2.4 We remark that when $\gamma_h > 0$ and $\eta\gamma_\varepsilon < 1$, it is a standard result in elliptic PDEs that (2.10) has a unique solution $u \in C^{3,v}(\mathbb{T}^d)$. We refer the reader to [21, 29] for more details.

When $\gamma_h = 0$ or $\gamma_h > 0$ is small, the continuum PDE (2.10) is dominated by the first-order terms and is better approximated by the first-order equation

$$u + \gamma_\varepsilon \rho^{-2} \operatorname{div}(\rho^2 bu) = \rho^{-1} v \quad \text{on } \mathbb{T}^d. \quad (2.14)$$

We state the first-order continuum limit as a separate result.

Theorem 2.5 (Convergence of the First-Order PageRank Problem) *Let $\rho \in C^{1,1}(\mathbb{T}^d)$, $b \in C^{1,1}(\mathbb{T}^d; \mathbb{R}^d)$ and $v \in C^{0,1}(\mathbb{T}^d)$. Assume that $\gamma_\varepsilon, \gamma_h \leq 1$, $\eta < 1$, and $\|Db\|_{L^\infty(\mathbb{T}^d)} \leq \frac{1}{2}(1 - \eta\gamma_\varepsilon)$. Let u_n be the solution to the PageRank problem (2.9) and let $u \in C^{0,1}(\mathbb{T}^d)$ be the viscosity solution of the PDE (2.14). Then there exists $C_1, C_2, c_1 > 0$ such that*

$$\max_{x \in \mathcal{X}_n} |u(x) - u_n(x)| \leq C_1 \sqrt{\lambda + \varepsilon + \gamma_h} \quad (2.15)$$

holds with probability at least $1 - C_2n \exp(-c_1nh^{d+2}\lambda^2) - C_2n \exp(-c_1nh^{d+2}(1 - \eta\gamma_\varepsilon)^2)$, where $0 < \lambda \leq 1$. We note that C_1 depends on $1 - \eta\gamma_\varepsilon$.

Remark 2.6 When $\eta\gamma_\varepsilon < 1$, it is a standard result in viscosity solution theory that (2.14) has a unique viscosity solution $u \in C(\mathbb{T}^d)$. We prove in Lemma 4.3 that when $\|Db\|_{L^\infty(\mathbb{T}^d)} \leq \frac{1}{2}(1 - \eta\gamma_\varepsilon)$ the viscosity solution u is Lipschitz continuous, so $u \in C^{0,1}(\mathbb{T}^d)$. We refer the reader to [18, 10] for more details on viscosity solutions.

We note that the continuum PDE (2.10) has reaction, advection and diffusion terms. The two reaction terms, u and $\rho^{-1}v$, are due to the teleportation step in PageRank. The term $\operatorname{div}(\rho^2 bu)$ is an advection term, which describes the advection of the quantity $\rho^2 u$ along the vector field b , and is due to the directional preference in the definition of the weights in a random directed geometric graph. Finally, the weighted diffusion term $\operatorname{div}(\rho^2 \nabla u)$ represents diffusion from the random walk step of PageRank.

Theorem 2.3 shows that if we scale $\varepsilon_n \sim \alpha_n$ and $h_n \sim \sqrt{\alpha_n}$, then the directional preference along the vector field b is balanced with the diffusion terms, and the limiting PDE (2.10) is second order. If $\varepsilon_n \ll \alpha_n$, then the directional preference is negligible in the limit, and the first-order terms in (2.10) disappear in the limit. Theorem 2.5 shows that if $h_n \ll \sqrt{\alpha_n}$, then the directional preference term dominates and the diffusion term is negligible, and the continuum PDE reduces to a first-order equation (2.14). If $\alpha_n \gg \max\{\varepsilon_n, h_n^2\}$, then the first- and second-order terms drop out of (2.10) and we simply get $u = \rho^{-1}v$. The intuition is that the teleportation happens too often and overwhelms the diffusion and directional preferences, yielding a trivial continuum limit.

Remark 2.7 We can analyse the characteristics of the first-order equation (2.14) to understand how the PageRank algorithm propagates information on the directed graph. The characteristic ordinary differential equations [21] are

$$\begin{cases} \dot{p}(s) = z(s) \nabla(\operatorname{div} b + 2 \nabla \log \rho \cdot b) \Big|_{x(s)} + D b(x(s)) p(s), \\ \dot{z}(s) = b(x(s)) \cdot p(s), \\ \dot{x}(s) = b(x(s)), \end{cases} \quad (2.16)$$

where $x(s)$ is the projected characteristic curve, $z(s) = u(x(s))$, and $p(s) = \nabla u(x(s))$. Hence, information is propagated along the integral curves of the vector field b , which represents the directional influence in the random directed geometric graph.

Remark 2.8 Theorems 2.3 and 2.5 are stated as finite sample size results, where n , ε , h , α and λ are fixed. If we consider the continuum limit as $n \rightarrow \infty$ and $\varepsilon_n, h_n, \alpha_n, \lambda_n \rightarrow 0$, then Theorems 2.3 and 2.5 inform us about how to relatively scale the parameters. We always assume $\varepsilon_n \leq \alpha_n$ and $h_n^2 \leq \alpha_n$, so that $\gamma_{\varepsilon_n}, \gamma_{h_n} \leq 1$. Thus, provided that

$$\lim_{n \rightarrow \infty} \frac{n h_n^{d+2} \lambda_n^2}{\log n} = \infty, \quad (2.17)$$

we may apply the Borel–Cantelli lemma to conclude that the rates hold almost surely as $n \rightarrow \infty$. This allows us to make a suitable choice for λ_n . Since the convergence rates scale with λ_n , we choose $\lambda_n \rightarrow 0$ as quickly as possible while ensuring that (2.17) holds. A reasonable choice is

$$\lambda_n = \frac{\log n}{\sqrt{n h_n^{d+2}}}. \quad (2.18)$$

With this choice of λ_n , we have the almost sure convergence rate of

$$\mathcal{O} \left(\frac{\log(n)^2}{\sqrt{n h_n^{d+2}}} + \varepsilon_n + h_n \right)$$

in Theorem 2.3, provided

$$\lim_{n \rightarrow \infty} \frac{n h_n^{d+2}}{\log(n)^2} = \infty. \quad (2.19)$$

The scaling in (2.19) is a standard scaling for pointwise consistency of graph Laplacians. Another way to phrase these results is to make the choice of

$$\lambda_n = \varepsilon_n + h_n$$

to match the terms in the error estimate in Theorem 2.3. In this case, we have an almost sure convergence rate of $\mathcal{O}(\varepsilon_n + h_n)$ provided

$$\lim_{n \rightarrow \infty} \frac{n h_n^{d+2} (\varepsilon_n + h_n)^2}{\log(n)} = \infty. \quad (2.20)$$

The same observations hold in the context of Theorem 2.5, except the rates are worse by a square root.

Remark 2.9 Theorems 2.3 and 2.5 can easily be rewritten in terms of the true PageRank vector $r_n(x)$. Due to Lemma 3.3, we have

$$\max_{x \in X_n} |\rho(x)u(x) - r_n(x)| \leq C_1(1 - \eta\gamma_\varepsilon)^{-1}(\lambda + \varepsilon + h)$$

in the context of Theorem 2.3, and

$$\max_{x \in X_n} |\rho(x)u(x) - r_n(x)| \leq C_1\sqrt{\lambda + \varepsilon + \gamma_h}$$

in the context of Theorem 2.5.

The PageRank vector r_n satisfies

$$\sum_{x \in X_n} r_n(x) = \sum_{x \in X_n} v(x),$$

as can be easily checked by summing both sides of (2.5). This forces r_n to be a probability distribution provided v is as well. The normalised PageRank vector u_n satisfies

$$\sum_{x \in X_n} \frac{d_n(x)}{nh^d} u_n(x) = \sum_{x \in X_n} v(x).$$

In the continuum, the solution u of (2.10) or (2.14) satisfies the analogous continuum version

$$\int_{\mathbb{T}^d} \rho^2 u \, dx = \int_{\mathbb{T}^d} \rho v \, dx,$$

which can be verified by multiplying both sides of (2.10) or (2.14) by ρ^2 and integrating by parts.

Remark 2.10 While the original PageRank problem is an eigenvector problem, the PageRank vector is an eigenvector of a probability transition matrix and not an eigenvector of a graph Laplacian. Thus, we cannot use the spectral properties of the graph Laplacian proven, for example, in [43, 25, 13] to address the eigenvalue problem (1.1). In fact, since the probability transition matrix becomes localised as $h, \varepsilon \rightarrow 0$, we lose the interpretation of PageRank as an eigenvector problem in the continuum. The localisation of the probability transition matrix is exactly what leads to an equation with a Laplacian in the continuum. A very simple analogue is the averaging operator

$$T_\varepsilon u(x) := \frac{1}{|B(x, \varepsilon)|} \int_{B(x, \varepsilon)} u(y) \, dy,$$

A function u satisfying $T_\varepsilon u = u$ is an eigenfunction of T_ε with eigenvalue $\lambda = 1$. In PDE theory, the equation $T_\varepsilon u = u$ is called the mean-value property and is satisfied by any harmonic function u . One can easily check that

$$\frac{1}{\varepsilon^2}(T_\varepsilon u(x) - u(x)) = C\Delta u(x) + O(\varepsilon)$$

for any smooth function u , where C depends only on d . Hence, as $\varepsilon \rightarrow 0$, eigenfunctions of T_ε are expected to converge to harmonic functions, which are solutions of Laplace's equation $\Delta u = 0$. The operator T_ε localises and becomes trivial as $\varepsilon \rightarrow 0$, since it reduces to pointwise evaluation

$T_0 u(x) = u(x)$. Thus, there is no meaningful way to think of harmonic functions as eigenfunctions of T_0 . An analogous, but more complicated, phenomenon occurs with the PageRank problem, as it also becomes localised in the continuum limit.

As an immediate application of Theorems 2.3 and 2.5, we prove asymptotic Lipschitz regularity of the PageRank vector, which shows that the ranking does not vary rapidly in feature space.

Corollary 2.11 (Lipschitz regularity) *Under the assumptions of Theorem 2.3, for $0 < \lambda \leq 1$ and with probability at least $1 - C_2 n \exp(-c_2 n h^{d+2} \lambda^2) - C_2 n \exp(-c_2 n h^{d+2} (1 - \eta \gamma_\varepsilon)^2)$ we have*

$$|u_n(x) - u_n(y)| \leq C|x - y| + C_1(1 - \eta \gamma_\varepsilon)^{-1}(\lambda + \varepsilon + h) \quad (2.21)$$

for all $x, y \in X_n$.

Proof By the triangle inequality

$$|u_n(x) - u_n(y)| \leq |u_n(x) - u(x)| + |u(x) - u(y)| + |u_n(y) - u(y)|.$$

We estimate the first and third term with Theorem 2.3, while the second is estimated by Lipschitzness of u . $\square$

Remark 2.12 Corollary 2.11 proves that u_n is approximately Lipschitz continuous, with jumps of size no larger than $O(\lambda + \varepsilon + h)$. We note that an analogous result to Corollary 2.11 can be stated under the assumptions of Theorem 2.5 as well.

We conclude this section by presenting an analogous continuum limit result for the evolution of the probability distribution of the random surfer $\mathbf{r}^k$. Similar to Eq. (1.1), the probability distribution $\mathbf{r}^k$ of the random surfer satisfies the evolution equation

$$\mathbf{r}^{k+1} = ((1 - \alpha)P + \alpha \mathbf{v} \mathbb{1}^T) \mathbf{r}^k. \quad (2.22)$$

Since $\mathbf{r}^k$ is a probability distribution, so $\mathbb{1}^T \mathbf{r}^k = 1$, we can also write the equation as

$$\mathbf{r}^{k+1} = (1 - \alpha)P \mathbf{r}^k + \alpha \mathbf{v}. \quad (2.23)$$

As before, we denote by $r_n(x, k)$ the x -component of $\mathbf{r}^k$, that is, $r_n(x, k)$ is the probability of finding the random surfer at vertex x after k steps on the random directed geometric graph of size n . Plugging this into (2.23), we find that $r_n(x, k)$ satisfies

$$r_n(x, k+1) = (1 - \alpha) \sum_{y \in X_n} \frac{\omega_n(y, x)}{d_n(y)} r_n(y, k) + \alpha v(x) \quad \text{for all } x \in X_n, \quad (2.24)$$

where $v(x)$ is the teleportation probability distribution. As before, we simplify the problem by defining the normalised distribution $u_n(x, k) := \frac{nh^d}{d_n(x)} r_n(x, k)$ and find that $u_n(x, k)$ satisfies

$$\frac{u_n(x, k+1) - u_n(x, k)}{\alpha} + u_n(x, k) - \gamma \mathcal{L}_n u_n(x, k) = \frac{nh^d}{d_n(x)} v(x) \quad \text{for all } x \in X_n, \quad (2.25)$$

For the initial condition, we take $u_n(x, 0) = g(x)$ for some given smooth function g . We can think of (2.25) as a discrete heat equation on the graph, describing the evolution of the normalised

distribution u_n of the random surfer. The stationary point of the evolution, as $k \rightarrow \infty$, is clearly the solution of the PageRank problem (2.9).

The continuum version of (2.25) is the reaction–advection–diffusion equation

$$\begin{cases} u_t + u + \gamma_\varepsilon \rho^{-2} \operatorname{div}(\rho^2 b u) - \frac{1}{2} \sigma_\Phi \gamma_h \rho^{-2} \operatorname{div}(\rho^2 \nabla u) = \rho^{-1} v, & \text{in } \mathbb{T}^d \times \{t > 0\} \\ u = g, & \text{on } \mathbb{T}^d \times \{t = 0\}. \end{cases} \quad (2.26)$$

This is verified by the following continuum limit result.

Theorem 2.13 (Continuum limit for random surfer) *Let $\rho \in C^{2,\nu}(\mathbb{T}^d)$, $b \in C^{2,\nu}(\mathbb{T}^d; \mathbb{R}^d)$, $v \in C^{1,\nu}(\mathbb{T}^d)$, and $g \in C^3(\mathbb{T}^d)$ for some $0 < \nu < 1$. Assume that $\gamma_\varepsilon \leq 1$, $0 < \gamma_h \leq 1$, and $\eta < 1$. Let $u_n(x, k)$ be the solution of (2.25) satisfying $u_n(x, 0) = g(x)$, and let $u \in C^3(\mathbb{T}^d)$ be the solution to the PDE (2.26). Then there exists $C_1, C_2, c_1, c_2 > 0$ with C_1 depending on $\gamma_h > 0$, such that when $\varepsilon + h \leq c_1(1 - \eta\gamma_\varepsilon)$ and $0 < \lambda \leq 1$, the event that*

$$\max_{x \in \mathcal{X}_n} |u(x, \alpha k) - u_n(x, k)| \leq C_1 \alpha k (\lambda + \varepsilon + h) \quad (2.27)$$

holds for all $k \geq 0$ has probability at least

$$1 - C_2 n \exp(-c_2 n h^{d+2} \lambda^2) - C_2 n \exp(-c_2 n h^{d+2} (1 - \eta\gamma_\varepsilon)^2).$$

Theorem 2.13 shows that the distribution of the random surfer can be approximated by the continuum PDE (2.26). The error estimates depend on λ, ε and h in a similar way as in Theorem 2.3. The main difference is the appearance of the term $k\alpha$, which corresponds to the time parameter in the continuum PDE (2.26), and is due to the accumulation of pointwise consistency errors over k steps.

Remark 2.14 *A first-order version of Theorem 2.13 can be proved, similar to Theorem 2.5. Since the statement and proof are very similar to Theorem 2.5, we omit the details.*

Remark 2.15 *We mention that another interesting perspective is the inverse problem of using graph-based numerical schemes, like the PageRank scheme, to numerically solve the continuum reaction–advection–diffusion equations. Modulo technical details, all of the results in this paper can be extended to the manifold setting, where $\mathbb{T}^d$ is replaced by a smooth compact and connected manifold $\mathcal{M}$ of dimension d embedded in $\mathbb{R}^D$ where $d < D$. Since graph-based numerical schemes learn the geometry of the manifold automatically, they may provide convenient numerical methods for solving PDEs on manifolds. We mention that ideas along these lines were mentioned in [8] for approximating the spectrum of the Laplace–Beltrami operator on a manifold of dimension higher than 2 or 3, where finite-element methods become cumbersome. This is an interesting direction to explore in future research.*

2.3 Outline

The paper will be organised as follows. In Section 3, we prove pointwise consistency with high probability for the PageRank operator $\mathcal{L}_n$ on a random directed geometric graph. This includes pointwise consistency to both first- and second-order continuum operators. In Section 4, we prove

our main results, Theorems 2.3, 2.5 and 2.13. Lastly, in Section 5, we present some numerical results to support our arguments.

3 Consistency for $\mathcal{L}_n$

In this section, we prove pointwise consistency for the operator $\mathcal{L}_n$ with both first- and second-order continuum operators. Throughout this section and the rest of the paper, we write $\omega_{xy} = \omega_n(x, y)$ and $d_x = d_n(x)$ for simplicity.

3.1 Concentration of measure and change of variables

We first recall a concentration inequality from [9].

Lemma 3.1 ([9, Remark 7]) *Let $Y_1, Y_2, \dots, Y_n$ be a sequence of i.i.d. random variables on $\mathbb{R}^d$ with density $f: \mathbb{R}^d \rightarrow \mathbb{R}$, let $\psi: \mathbb{R}^d \rightarrow \mathbb{R}$ be bounded and Borel measurable with compact support in a bounded open set $\Omega \subset \mathbb{R}^d$, and define*

$$Y = \sum_{i=1}^n \psi(Y_i).$$

Then for any $0 \leq \lambda \leq 1$,

$$\mathbb{P}\left[|Y - \mathbb{E}(Y)| > \|f\|_\infty \|\psi\|_\infty n |\Omega| \lambda\right] \leq 2 \exp\left(-\frac{1}{4} \|f\|_\infty n |\Omega| \lambda^2\right), \quad (3.1)$$

where $|\Omega|$ denotes the Lebesgue measure of Ω .

When we apply the lemma to the degree d_y , we need to compute the expected value of d_y , which is the integral

$$\mathbb{E}(d_y) = n \int_{E_y} \Phi\left(\frac{|x - y - \varepsilon b(y)|}{h}\right) \rho(x) dx.$$

For this computation, and others, we require asymptotic expansions in the change of variables formulas, which is provided by the following result.

Lemma 3.2 (Change of Variables) *Let $g: \mathbb{R}^d \rightarrow \mathbb{R}$ be continuous and assume b is C^1 . Then using the change of variables $z = \frac{x - y - \varepsilon b(y)}{h}$, we have*

$$\int_{\mathbb{R}^d} \Phi\left(\frac{|x - y - \varepsilon b(y)|}{h}\right) g(x) dx = h^d \int_{\mathbb{R}^d} \Phi(|z|) g(y + hz + \varepsilon b(y)) dz, \quad (3.2)$$

and

$$\begin{aligned} & \int_{\mathbb{R}^d} \Phi\left(\frac{|x - y - \varepsilon b(y)|}{h}\right) g(y) dy \\ &= h^d \int_{\mathbb{R}^d} \Phi(|z|) g(x - hz - \varepsilon b(x) + \mathcal{O}(\varepsilon h + \varepsilon^2)) (1 - \varepsilon \operatorname{div} b(x) + \mathcal{O}(\varepsilon h + \varepsilon^2)) dz, \end{aligned} \quad (3.3)$$

for sufficiently small $\varepsilon > 0$.

Proof The proof is a straightforward change of variables when we integrate over x in (3.2). In the case where we integrate over y in (3.3), we define the change of variables function $G: \mathbb{R}^d \rightarrow \mathbb{R}^d$ by

$$G(y) = \frac{x - y - \varepsilon b(y)}{h},$$

where the vector field b is extended periodically from $\mathbb{T}^d$ to $\mathbb{R}^d$. For ε sufficiently small, G is a C^1 diffeomorphism of $\mathbb{R}^d$. Indeed, to see that G is bijective, first note that the problem of solving $G(y) = z$ for y , given z , is equivalent to the fixed point problem

$$y = \Psi(y) := x - zh - \varepsilon b(y).$$

The mapping $\Psi: \mathbb{R}^d \rightarrow \mathbb{R}^d$ is a contraction, provided $\varepsilon < \|b\|_{C^{0,1}}^{-1}$, and so the invertibility of G follows from Banach's fixed point theorem. For $\varepsilon < \|b\|_{C^{0,1}}^{-1}$ the Jacobian matrix $D_y G(y) = -h^{-1}(I + \varepsilon D_y b(y))$ is invertible, and so the smoothness of G^{-1} follows from the inverse function theorem.

We now use the change of variables $z = G(y)$ in (3.3) to obtain

$$\int_{\mathbb{R}^d} \Phi\left(\frac{|x - y - \varepsilon b(y)|}{h}\right) g(y) dy = \int_{\mathbb{R}^d} \Phi(|z|) g(G^{-1}(z)) |\det D_y G(G^{-1}(z))|^{-1} dz. \quad (3.4)$$

In the rest of the proof, we write $y = G^{-1}(z)$ for convenience. Then we have

$$|\det D_y G(y)|^{-1} = h^d |\det(I + \varepsilon D b(y))|^{-1}. \quad (3.5)$$

We now use the Taylor expansion

$$\det(I + \varepsilon A) = 1 + \varepsilon \text{Tr}(A) + \mathcal{O}(\varepsilon^2)$$

for $A = D b(y)$ to obtain

$$\det(I + \varepsilon D b(y)) = 1 + \varepsilon \text{Tr}(D b(y)) + \mathcal{O}(\varepsilon^2) = 1 + \varepsilon \text{div} b(x) + \mathcal{O}(\varepsilon h + \varepsilon^2).$$

Thus, for sufficiently small $\varepsilon > 0$ we have

$$|\det D_y G(y)|^{-1} = h^d (1 - \varepsilon \text{div} b(x) + \mathcal{O}(\varepsilon h + \varepsilon^2)). \quad (3.6)$$

Since $\Phi(t) = 0$ for $t > 2$ we have that $|x - y| \leq 2h + \varepsilon |b(y)| \leq C(h + \varepsilon)$, and so

$$G^{-1}(z) = y = x - hz - \varepsilon b(y) = x - hz - \varepsilon b(x) + \mathcal{O}(\varepsilon h + \varepsilon^2). \quad (3.7)$$

Substituting (3.6) and (3.7) into (3.4) completes the proof. $\square$

3.2 Pointwise consistency

We now turn to the main pointwise consistency results. We begin with a standard result for the degree.

Lemma 3.3 (Asymptotics for the degree) *For any $0 < \lambda \leq 1$, the degree term satisfies*

$$d_y = nh^d \rho(y) + \mathcal{O}(nh^d(\lambda + \varepsilon + h^2))$$

with probability at least $1 - 2 \exp(-c\lambda^2 nh^d)$, where c is a constant independent of n .

Proof We use part (i) in Proposition 3.2 to compute

$$\begin{aligned}\mathbb{E}(d_y) &= \mathbb{E}\left(\sum_x \omega_{yx}\right) = n \int_{E_y} \Phi\left(\frac{|x-y-\varepsilon b(y)|}{h}\right) \rho(x) dx \\ &= nh^d \int_{B(0,2)} \Phi(|z|) \rho(y+hz+\varepsilon b(y)) dz \\ &= nh^d (\rho(y) + \mathcal{O}(\varepsilon + h^2)).\end{aligned}$$

By Lemma 3.1, we see that

$$\mathbb{P}\left[\left|d_y - \mathbb{E}[d_y]\right| > C\lambda nh^d\right] \leq 2 \exp\left(-\frac{1}{4}cnh^d\lambda^2\right)$$

for any $0 \leq \lambda \leq 1$. Combining the observations above completes the proof. $\square$

We now state and prove the consistency result.

Theorem 3.4 (Consistency for the Second-Order PageRank Operator) *There exists constants $C, c > 0$ such that for any $0 < \lambda \leq 1$ the event that*

$$\begin{aligned}\frac{d_x}{\rho(x)nh^d} \mathcal{L}_n \varphi(x) &= -\rho^{-2} \operatorname{div}(\rho^2 b \varphi) \varepsilon + \frac{\sigma_\Phi}{2} \rho^{-2} \operatorname{div}(\rho^2 \nabla \varphi) h^2 \Big|_x \\ &\quad + \mathcal{O}\left((\lambda h + \lambda h^{-1} \varepsilon + \varepsilon^2 + h^3 + \varepsilon h) \|\varphi\|_{C^{2,1}(\mathbb{T}^d)}\right)\end{aligned}\tag{3.8}$$

holds for all $\varphi \in C^{2,1}(\mathbb{T}^d)$ and $x \in X_n$ has probability at least $1 - Cn \exp(-cnh^d\lambda^2)$.

Proof Fix $x \in \mathbb{T}^d$, take $\varphi \in C^{2,1}(\mathbb{R}^d)$ to be a test function and let $p = D\varphi(x)$ and $a_{ij} = \varphi_{x_i x_j}(x)$. We apply the operator $\mathcal{L}_n$ to φ at x and take a second-order Taylor expansion at x of the $\varphi(y)$ inside the summation, which gives us

$$\begin{aligned}d_x \mathcal{L}_n \varphi(x) &= \sum_{y \in X_n \setminus \{x\}} (\omega_{yx} \varphi(y) - \omega_{xy} \varphi(x)) \\ &= \sum_{y \in X_n \setminus \{x\}} (\omega_{yx} - \omega_{xy}) \varphi(x) + \sum_{i=1}^d p_i \sum_{y \in X_n \setminus \{x\}} \omega_{yx} (y_i - x_i) \\ &\quad + \frac{1}{2} \sum_{i,j=1}^d a_{ij} \sum_{y \in X_n \setminus \{x\}} \omega_{yx} (y_i - x_i)(y_j - x_j) + \mathcal{O}\left((\varepsilon^3 + h^3) \beta \sum_{y \in X_n \setminus \{x\}} \omega_{yx}\right),\end{aligned}$$

where $\beta = \|\varphi\|_{C^{2,1}(\mathbb{T}^d)}$, and the $\varepsilon^3 + h^3$ in the remainder term comes from the scaling of $|y - x|$. Since $\varphi(x)$ and its derivatives are factored out from the summations, the probability estimates that follow are independent of φ and hold uniformly over all smooth test functions.

Noting that $|\omega_{xy} - \omega_{yx}| \leq Ch^{-1}\varepsilon$, we have by Lemma 3.1 that each of

$$\begin{aligned} \sum_{y \in X_n \setminus \{x\}} \omega_{yx} &\leq Cnh^d, \\ \frac{1}{n} \sum_{y \in X_n \setminus \{x\}} (\omega_{yx} - \omega_{xy}) &= \int_{\mathbb{T}^d} (\omega_{yx} - \omega_{xy}) \rho(y) dy + \mathcal{O}(\lambda h^{-1}\varepsilon), \\ \frac{1}{n} \sum_{y \in X_n \setminus \{x\}} \omega_{yx}(y_i - x_i) &= \int_{\mathbb{T}^d} \omega_{yx}(y_i - x_i) \rho(y) dy + \mathcal{O}(\lambda(\varepsilon + h)), \text{ and} \\ \frac{1}{n} \sum_{y \in X_n \setminus \{x\}} \omega_{yx}(y_i - x_i)(y_j - x_j) &= \int_{\mathbb{T}^d} \omega_{yx}(y_i - x_i)(y_j - x_j) \rho(y) dy + \mathcal{O}(\lambda(\varepsilon^2 + h^2)) \end{aligned}$$

holds with probability at least $1 - 2 \exp(-cnh^d\lambda^2)$ for any $0 < \lambda \leq 1$. Combining the observations above we have with probability at least $1 - C \exp(-cnh^d\lambda^2)$ that

$$\begin{aligned} \frac{d_x}{nh^d} \mathcal{L}_n \varphi(x) &= \frac{1}{h^d} \int_{\mathbb{T}^d} (\omega_{yx} - \omega_{xy}) \rho(y) dy \varphi(x) + \sum_{i=1}^d \frac{p_i}{h^d} \int_{\mathbb{T}^d} \omega_{yx}(y_i - x_i) \rho(y) dy \\ &\quad + \frac{1}{2} \sum_{i,j=1}^d \frac{a_{ij}}{h^d} \int_{\mathbb{T}^d} \omega_{yx}(y_i - x_i)(y_j - x_j) \rho(y) dy \\ &\quad + \mathcal{O}(\lambda(h + h^{-1}\varepsilon)\beta + (\varepsilon^3 + h^3)\beta). \end{aligned} \quad (3.9)$$

We now compute asymptotic expansions for all the terms in (3.9). By Lemma 3.2, we have

$$\begin{aligned} \frac{1}{h^d} \int_{\mathbb{T}^d} \omega_{yx} \rho(y) dy &= \frac{1}{h^d} \int_{\mathbb{T}^d} \Phi\left(\frac{|x - y - \varepsilon b(y)|}{h}\right) \rho(y) dy \\ &= \underbrace{\int_{B(0,2)} \Phi(|z|) \rho(x - hz - \varepsilon b(x) + \mathcal{O}(\varepsilon h + \varepsilon^2)) dz}_{A} (1 - \varepsilon \operatorname{div} b(x) + \mathcal{O}(\varepsilon h + \varepsilon^2)). \end{aligned}$$

We now compute

$$\begin{aligned} A &= \int_{B(0,2)} \Phi(|z|) \left(\rho(x) - \nabla \rho(x) \cdot (hz + \varepsilon b(x)) + \frac{h^2}{2} z^T \nabla^2 \rho(x) z + \mathcal{O}(\varepsilon h + \varepsilon^2) \right) dz \\ &= \rho(x) - \nabla \rho(x) \cdot b(x) \varepsilon + \frac{\sigma_\Phi}{2} \Delta \rho(x) h^2 + \mathcal{O}(\varepsilon h + \varepsilon^2). \end{aligned}$$

Therefore,

$$\begin{aligned} \frac{1}{h^d} \int_{\mathbb{T}^d} \omega_{yx} \rho(y) dy &= \left(\rho(x) - \nabla \rho(x) \cdot b(x) \varepsilon + \frac{\sigma_\Phi}{2} \Delta \rho(x) h^2 + \mathcal{O}(\varepsilon h + \varepsilon^2) \right) (1 - \varepsilon \operatorname{div} b(x) + \mathcal{O}(\varepsilon h + \varepsilon^2)) \\ &= \rho(x) - \nabla \rho(x) \cdot b(x) \varepsilon + \frac{\sigma_\Phi}{2} \Delta \rho(x) h^2 - \rho(x) \operatorname{div} b(x) \varepsilon + \mathcal{O}(\varepsilon h + \varepsilon^2). \end{aligned}$$

By Lemma 3.2, we have

$$\begin{aligned}
 & \frac{1}{h^d} \int_{\mathbb{T}^d} \omega_{xy} \rho(y) dy \\
 &= \frac{1}{h^d} \int_{\mathbb{T}^d} \Phi \left(\frac{|y - x - \varepsilon b(x)|}{h} \right) \rho(y) dy \\
 &= \int_{B(0,2)} \Phi(|z|) \rho(x + hz + \varepsilon b(x)) dz \\
 &= \int_{B(0,2)} \Phi(|z|) \left(\rho(x) + \nabla \rho(x) \cdot (hz + \varepsilon b(x)) + \frac{h^2}{2} z^T \nabla^2 \rho(x) z + \mathcal{O}(\varepsilon h + \varepsilon^2) \right) dz \\
 &= \rho(x) + \nabla \rho(x) \cdot b(x) \varepsilon + \frac{\sigma_\Phi}{2} \Delta \rho(x) h^2 + \mathcal{O}(\varepsilon h + \varepsilon^2).
 \end{aligned}$$

Therefore,

$$\frac{1}{h^d} \int_{\mathbb{T}^d} (\omega_{yx} - \omega_{xy}) \rho(y) dy = -2 \nabla \rho(x) \cdot b(x) \varepsilon - \rho(x) \operatorname{div} b(x) \varepsilon + \mathcal{O}(\varepsilon h + \varepsilon^2). \quad (3.10)$$

By Lemma 3.2 again, we have

$$\begin{aligned}
 & \frac{1}{h^d} \int_{\mathbb{T}^d} \omega_{yx} (y_i - x_i) \rho(y) dy \\
 &= \frac{1}{h^d} \int_{\mathbb{T}^d} \Phi \left(\frac{|x - y - \varepsilon b(y)|}{h} \right) (y_i - x_i) \rho(y) dy \\
 &= - \int_{B(0,2)} \Phi(|z|) \rho(x - hz + \mathcal{O}(\varepsilon)) (z_i h + b_i(x) \varepsilon + \mathcal{O}(\varepsilon h + \varepsilon^2)) dz (1 + \mathcal{O}(\varepsilon)) \\
 &= - \int_{B(0,2)} \Phi(|z|) (\rho(x) - \nabla \rho(x) \cdot zh + \mathcal{O}(\varepsilon + h^2)) (z_i h + b_i(x) \varepsilon + \mathcal{O}(\varepsilon h + \varepsilon^2)) dz \\
 &= - \int_{B(0,2)} \Phi(|z|) (\rho(x) z_i h + \rho(x) b_i(x) \varepsilon - (\nabla \rho(x) \cdot z) z_i h^2 + \mathcal{O}(\varepsilon h + h^3 + \varepsilon^2)) dz \\
 &= -\rho(x) b_i(x) \varepsilon + \sigma_\Phi \rho_{x_i}(x) h^2 + \mathcal{O}(\varepsilon h + h^3 + \varepsilon^2).
 \end{aligned} \quad (3.11)$$

Finally, another application of Lemma 3.2 yields

$$\begin{aligned}
 & \frac{1}{h^d} \int_{\mathbb{T}^d} \omega_{yx} (y_i - x_i) (y_j - x_j) \rho(y) dy \\
 &= \frac{1}{h^d} \int_{\mathbb{T}^d} \Phi \left(\frac{|x - y - \varepsilon b(y)|}{h} \right) (y_i - x_i) (y_j - x_j) \rho(y) dy \\
 &= \int_{B(0,2)} \Phi(|z|) (\rho(x) + \mathcal{O}(h + \varepsilon)) (z_i z_j h^2 + \mathcal{O}(\varepsilon h + \varepsilon^2)) dz \\
 &= \sigma_\Phi \rho(x) \delta_{ij} h^2 + \mathcal{O}(h^3 + \varepsilon h + \varepsilon^2),
 \end{aligned} \quad (3.12)$$

where $\delta_{ij} = 1$ if $i = j$ and $\delta_{ij} = 0$ otherwise. Combining (3.10), (3.11) and (3.12) with (3.9), we have

$$\begin{aligned} \frac{d_x}{nh^d} \mathcal{L}_n \varphi(x) &= -(2\nabla \rho(x) \cdot b(x) + \rho(x) \operatorname{div} b(x)) \varphi(x) \varepsilon \\ &\quad + \nabla \varphi(x) \cdot (\sigma_\Phi \nabla \rho(x) h^2 - \rho(x) b(x) \varepsilon) + \frac{\sigma_\Phi}{2} \rho(x) \Delta \varphi(x) h^2 \\ &\quad + \mathcal{O}(\lambda(h + h^{-1} \varepsilon) \beta + (\varepsilon^2 + h^3 + \varepsilon h) \beta). \end{aligned} \quad (3.13)$$

We now divide both sides of (3.13) by $\rho(x)$ and use the identities

$$\rho^{-2} \operatorname{div}(\rho^2 b \varphi) = \varphi(\operatorname{div} b + 2\nabla \log \rho \cdot b) + \nabla \varphi \cdot b,$$

and

$$\frac{1}{2} \rho^{-2} \operatorname{div}(\rho^2 \nabla \varphi) = \frac{1}{2} \Delta \varphi + \nabla \log \rho \cdot \nabla \varphi.$$

to establish that the event that (3.8) holds for all $\varphi \in C^{2,1}(\mathbb{T}^d)$ and a fixed $x \in \mathbb{T}^d$ has probability at least $1 - C \exp(-c n h^d \lambda^2)$.

To establish that (3.8) holds for all $x \in X_n$, we first condition on the random variable x_1 . The remaining points $X_n \setminus \{x_1\} = \{x_2, x_3, \dots, x_{n-1}\}$ are an independent *i.i.d.* sample of size $n - 1$, and we can write, as at the start of the proof, that

$$d_{x_1} \mathcal{L}_n \varphi(x_1) = \sum_{y \in X_n \setminus \{x_1\}} (\omega_{yx_1} \varphi(y) - \omega_{x_1 y} \varphi(x_1)).$$

Since the sum on the right is over $n - 1$ *i.i.d.* random variables, we can apply the same argument as above, and the law of conditional probability, to establish that the event that (3.8) holds for $x = x_1$ and all $\varphi \in C^{2,1}(\mathbb{T}^d)$ has probability at least $1 - C \exp(-c(n - 1)h^d \lambda^2)$. We then repeat the argument, conditioning on each of $x_2, x_3, \dots, x_n$, and union bound over all n points to obtain that (3.8) holds for all $x \in X_n$ and $\varphi \in C^{2,1}(\mathbb{T}^d)$ with probability at least $1 - Cn \exp(-c(n - 1)h^d \lambda^2)$. The proof is completed by using that $n - 1 \geq n/2$ for $n \geq 2$ to simplify the probability. $\square$

We also have a corresponding consistency result when the continuum PDE is first order.

Theorem 3.5 (Consistency for the First-Order PageRank Operator) *There exists constants $C, c > 0$ such that for any $0 < \lambda \leq 1$ the event that*

$$\frac{d_x}{\rho(x) n h^d} \mathcal{L}_n \varphi(x) = -\rho^{-2} \operatorname{div}(\rho^2 b \varphi) \varepsilon + \mathcal{O}((\lambda h + \lambda h^{-1} \varepsilon + \varepsilon^2 + h^2) \|\varphi\|_{C^{1,1}(\mathbb{T}^d)}) \quad (3.14)$$

holds for all $\varphi \in C^{1,1}(\mathbb{T}^d)$ and $x \in X_n$ has probability at least $1 - Cn \exp(-c n h^d \lambda^2)$.

Proof The proof follows closely to that of Theorem 3.4, so we sketch it here. Fix $x \in \mathbb{T}^d$, take $\varphi \in C^{1,1}(\mathbb{R}^d)$ to be a test function and let $p = D\varphi(x)$. We have

$$\begin{aligned} d_x \mathcal{L}_n \varphi(x) &= \sum_{y \in X_n} \omega_{yx} \varphi(y) - d_x \varphi(x) \\ &= \sum_{y \in X_n} (\omega_{yx} - \omega_{xy}) \varphi(x) + \sum_{i=1}^d p_i \sum_{y \in X_n} \omega_{yx} (y_i - x_i) + \mathcal{O}\left((\varepsilon^2 + h^2) \beta \sum_y \omega_{yx}\right), \end{aligned}$$

where $\beta = \|\varphi\|_{C^{1,1}(\mathbb{T}^d)}$. Thus, with probability at least $1 - C \exp(-cnh^d \lambda^2)$ we have that

$$\begin{aligned} \frac{d_x}{nh^d} \mathcal{L}_n \varphi(x) &= \frac{1}{h^d} \int_{\mathbb{T}^d} (\omega_{yx} - \omega_{xy}) \rho(y) dy \varphi(x) + \sum_{i=1}^d \frac{p_i}{h^d} \int_{\mathbb{T}^d} \omega_{yx} (y_i - x_i) \rho(y) dy \\ &\quad + \mathcal{O}(\lambda(h + h^{-1}\varepsilon)\beta + (\varepsilon^2 + h^2)\beta). \end{aligned} \quad (3.15)$$

By (3.10) and (3.11), we have

$$\begin{aligned} \frac{d_x}{nh^d} \mathcal{L}_n \varphi(x) &= -(2\nabla \rho(x) \cdot b(x) + \rho(x) \operatorname{div} b(x)) \varphi(x) \varepsilon - \nabla \varphi(x) \cdot \rho(x) b(x) \varepsilon \\ &\quad + \mathcal{O}(\lambda(h + h^{-1}\varepsilon)\beta + (\varepsilon^2 + h^2)\beta). \end{aligned}$$

Divide both sides by $\rho(x)$ and use the identity

$$\rho^{-2} \operatorname{div}(\rho^2 b \varphi) = \varphi(\operatorname{div} b + 2\nabla \log \rho \cdot b) + \nabla \varphi \cdot b$$

to complete the proof. $\square$

4 Convergence proofs

We now prove our main results. We first need a stability estimate for the PageRank problem (2.9).

Lemma 4.1 (ℓ_∞ Stability for the PageRank Operator) *Assume that $\gamma_\varepsilon, \gamma_h \leq 1$ and $\eta < 1$, where η is defined in (2.12) and $\gamma_\varepsilon, \gamma_h$ in (2.11). There exists $C, K, c > 0$ such that with probability at least $1 - Cn \exp(-cnh^{d+2}(1 - \eta\gamma_\varepsilon)^2)$, if $\varepsilon + h \leq K(1 - \eta\gamma_\varepsilon)$ and $u, v : X_n \rightarrow \mathbb{R}$ satisfy*

$$u(x) - \gamma \mathcal{L}_n u(x) = \frac{nh^d}{d_x} v(x) \quad \text{for all } x \in X_n \quad (4.1)$$

with $\gamma = (1 - \alpha)/\alpha$, then it holds that

$$\max_{x \in X_n} |u(x)| \leq 2(1 - \eta\gamma_\varepsilon)^{-1} \max_{x \in X_n} |\rho(x)^{-1} v(x)|. \quad (4.2)$$

Proof We use a maximum principle argument. Let $x_0 \in X_n$ be a point where u attains its maximum value. Setting $\varphi \equiv 1$, we have

$$\mathcal{L}_n u(x_0) = \frac{1}{d_{x_0}} \sum_{y \in X_n} \omega_{yx_0} u(y) - u(x_0) \leq \frac{1}{d_{x_0}} \sum_{y \in X_n} \omega_{yx_0} u(x_0) - u(x_0) = u(x_0) \mathcal{L}_n \varphi(x_0).$$

It follows that

$$\frac{d_{x_0}}{nh^d} (1 - \gamma \mathcal{L}_n \varphi(x_0)) u(x_0) \leq v(x_0). \quad (4.3)$$

By Theorem 3.4, we have

$$\frac{d_{x_0}}{\rho(x) nh^d} \mathcal{L}_n \varphi(x_0) = -\rho^{-2} \operatorname{div}(\rho^2 b) \varepsilon + \mathcal{O}(\lambda h + \lambda h^{-1} \varepsilon + \varepsilon^2 + h^3 + \varepsilon h),$$

with probability at least $1 - Cn \exp(-cnh^d \lambda^2)$. In particular, observe that $\operatorname{div}(\rho^2 \nabla \varphi) = 0$ in Theorem 3.4, since $\varphi \equiv 1$ is constant. Setting $\lambda = \delta h$ for $0 < \delta \leq h^{-1}$ and recalling (2.12), we have

$$\frac{d_{x_0}}{\rho(x)nh^d} \gamma |\mathcal{L}_n \varphi(x)| \leq \eta \gamma_\varepsilon + C(\gamma_h + \gamma_\varepsilon)(\delta + h + \varepsilon) \leq \eta \gamma_\varepsilon + C(\delta + h + \varepsilon),$$

with probability at least $1 - Cn \exp(-cnh^{d+2}\delta^2)$. By Lemma 3.3, we have

$$\frac{d_{x_0}}{\rho(x)nh^d} = 1 + \mathcal{O}(\varepsilon + h)$$

with probability at least $1 - 2n \exp(-cnh^{d+2})$. Inserting these observations into (4.3), we have

$$\rho(x_0)(1 - \eta \gamma_\varepsilon - C(\delta + h + \varepsilon))u(x_0) \leq v(x_0),$$

for a constant $C > 0$. Hence, selecting $\delta = (1 - \eta \gamma_\varepsilon)/(4C)$ and restricting $h + \varepsilon \leq (1 - \eta \gamma_\varepsilon)/(4C)$, we have

$$\frac{1}{2} \rho(x_0)(1 - \eta \gamma_\varepsilon)u(x_0) \leq v(x_0),$$

with probability at least $1 - Cn \exp(-cnh^{d+2}(1 - \eta \gamma_\varepsilon)^2)$. Therefore,

$$\max_{x \in X_n} u(x) \leq 2(1 - \eta \gamma_\varepsilon)^{-1} \max_{x \in X_n} |\rho(x)^{-1} v(x)|$$

holds with probability at least $1 - Cn \exp(-cnh^{d+2}(1 - \eta \gamma_\varepsilon)^2)$ provided $\varepsilon + h \leq K(1 - \eta \gamma_\varepsilon)$.

For the proof of the other direction, we set $\bar{u}(x) = -u(x)$, and note that $\bar{u}$ satisfies

$$\bar{u}(x) - \gamma \mathcal{L}_n \bar{u}(x) = -\frac{nh^d}{d_x} v(x) \quad \text{for all } x \in X_n.$$

The argument in the first part of the proof yields

$$\max_{x \in X_n} \bar{u}(x) \leq 2(1 - \eta \gamma_\varepsilon)^{-1} \max_{x \in X_n} |\rho(x)^{-1} v(x)|,$$

which completes the proof. $\square$

Given the stability estimate from Lemma 4.1, we can now prove Theorem 2.3.

Proof of Theorem 2.3 Given the assumptions on ρ , b and v , the solution u of the continuum PDE (2.10) belongs to $C^3(\mathbb{T}^d)$, and $\|u\|_{C^3(\mathbb{T}^d)}$ depends on the ellipticity constant of the equation $\sigma_\Phi \gamma_h$ [29, 21]. Thus, applying Lemma 3.3, and Theorem 3.4 with $\lambda = \delta h$ yields

$$\begin{aligned} & \frac{d_x}{\rho(x)nh^d} u(x) - \frac{d_x}{\rho(x)nh^d} \gamma \mathcal{L}_n u(x) \\ &= u(x) + \gamma_\varepsilon \rho^{-2} \operatorname{div}(\rho^2 b u) - \frac{1}{2} \sigma_\Phi \gamma_h \rho^{-2} \operatorname{div}(\rho^2 \nabla u) + \mathcal{O}(\delta + \varepsilon + h) \\ &= \frac{v(x)}{\rho(x)} + \mathcal{O}(\delta + \varepsilon + h) \end{aligned}$$

for all $x \in X_n$ with probability at least $1 - Cn \exp(-cnh^{d+2}\delta^2)$ for any $0 < \delta \leq h^{-1}$, where we used that $\gamma_\varepsilon, \gamma_h \leq 1$. Therefore,

$$u(x) - \gamma \mathcal{L}_n u(x) = \frac{nh^d}{d_x} (v(x) + \mathcal{O}(\delta + \varepsilon + h))$$

for all $x \in X_n$ with probability at least $1 - Cn \exp(-cnh^{d+2}\delta^2)$.

Consider now the difference $w_n(x) = u(x) - u_n(x)$, where u_n solves the PageRank problem (2.9). Then w_n satisfies

$$w_n(x) - \gamma \mathcal{L}_n w_n(x) = \frac{nh^d}{d_x} \mathcal{O}(\delta + \varepsilon + h)$$

for all $x \in X_n$ with probability at least $1 - Cn \exp(-cnh^{d+2}\delta^2)$. Applying Lemma 4.1 completes the proof. $\square$

We now give the proof of Theorem 2.13.

Proof of Theorem 2.13 The proof is similar to the proof of Theorem 2.3, so we sketch the outline, omitting some details.

First, note that

$$\frac{u(x, \alpha k + \alpha) - u(x, \alpha k)}{\alpha} = u_t(x, \alpha k) + \mathcal{O}(\alpha).$$

Proceeding now as in the proof of Theorem 2.3, we use Theorem 3.4, Lemma 3.3, and the observation above to deduce that

$$\frac{u(x, \alpha k + \alpha) - u(x, \alpha k)}{\alpha} + u(x, \alpha k) - \gamma \mathcal{L}_n u(x, \alpha k) = \frac{nh^d}{d_n(x)} v(x) + O(\lambda + \varepsilon + h), \quad (4.4)$$

with probability at least $1 - Cn \exp(-cnh^{d+2}\lambda^2)$ for $0 < \lambda \leq 1$, where we also used that $\gamma_\varepsilon \leq 1$. Let $\varphi(x) = 1$ for all $x \in X_n$. As in the proof of Lemma 4.1 we have that

$$(1 - \alpha)|\mathcal{L}_n \varphi(x)| \leq \alpha \eta \gamma_\varepsilon + C\alpha(\delta + h + \varepsilon) \quad (4.5)$$

with probability at least $1 - Cn \exp(-cnh^{d+2}\delta^2)$, where $0 < \delta \leq 1$ will be chosen later. For the rest of the proof, we assume the events above hold true.

Define $w_n(x, k) = u(x, \alpha k) - u_n(x, k)$. We claim that

$$w_n(x, k) \leq Ck\alpha(\lambda + \varepsilon + h) =: M_k. \quad (4.6)$$

The proof of the other direction is similar, and this will complete the proof. To prove (4.6), we use a comparison principle argument that proceeds by induction. The base case $k = 0$ is trivial, since $w_n(x, 0) = 0$. Assume that (4.6) is true for some $k \geq 0$. Then subtracting (4.4) and (2.25), we find that w_n satisfies

$$\frac{w_n(x, k+1) - w_n(x, k)}{\alpha} + w_n(x, k) - \gamma \mathcal{L}_n w_n(x, k) \leq C(\lambda + \varepsilon + h)$$

for all $x \in X_n$ and $k \geq 0$. Rearranging this, we obtain

$$\begin{aligned} w_n(x, k+1) &\leq (1 - \alpha)(w_n(x, k) + \mathcal{L}_n w_n(x, k)) + C\alpha(\lambda + \varepsilon + h) \\ &= (1 - \alpha) \frac{1}{d_x} \sum_{y \in X_n} w_{yx} w_n(y, k) + C\alpha(\lambda + \varepsilon + h) \\ &\leq (1 - \alpha) \frac{M_k}{d_x} \sum_{y \in X_n} w_{yx} + C\alpha(\lambda + \varepsilon + h) \\ &= M_k(1 - \alpha)(1 + \mathcal{L}_n \varphi(x)) + C\alpha(\lambda + \varepsilon + h), \end{aligned}$$

since $w_n(x, k) \leq M_k$ for all $x \in X_n$. Applying (4.5), we have

$$w_n(x, k+1) \leq M_k (1 - (1 - \eta\gamma_\varepsilon)\alpha + C\alpha(\delta + h + \varepsilon)) + C\alpha(\lambda + \varepsilon + h).$$

Hence, choosing $\delta = (1 - \eta\gamma_\varepsilon)/(4C)$ and restricting $h + \varepsilon \leq (1 - \eta\gamma_\varepsilon)/(4C)$, we have

$$w_n(x, k+1) \leq M_k + C\alpha(\lambda + \varepsilon + h) = M_{k+1}.$$

The claim (4.6) is thus established by induction, and this completes the proof. $\square$

We now turn our attention to proving the first-order rate, which hinges on the following observation that our scheme is *monotone*.

Proposition 4.2 (Monotonicity) *Let $u, v : X_n \rightarrow \mathbb{R}$ and $x_0 \in X_n$ such that $u(x_0) = v(x_0)$ and $u \leq v$. Then $\mathcal{L}_n u(x_0) \leq \mathcal{L}_n v(x_0)$.*

Proof The proof is immediate, since

$$\mathcal{L}_n u(x_0) = \frac{1}{d_x} \sum_{y \in X_n} \omega_{yx} u(y) - u(x_0) \leq \frac{1}{d_x} \sum_{y \in X_n} \omega_{yx} v(y) - v(x_0) = \mathcal{L}_n v(x_0). \quad \square$$

In order to prove a convergence rate for the first-order continuum limit, we require a Lipschitz estimate on the viscosity solution u of (2.14). The result follows a standard maximum principle argument, which we include for completeness.

Lemma 4.3 *Let $\rho \in C^{1,1}(\mathbb{T}^d)$, $b \in C^{1,1}(\mathbb{T}^d; \mathbb{R}^d)$ and $v \in C^{0,1}(\mathbb{T}^d)$. Assume that $\gamma_\varepsilon \leq 1$, and $\eta < 1$. Let $u \in C(\mathbb{T}^d)$ be the viscosity solution of the PDE (2.14). If $\|Db\|_{L^\infty(\mathbb{T}^d)} \leq \frac{1}{2}(1 - \eta\gamma_\varepsilon)$, then $u \in C^{0,1}(\mathbb{T}^d)$ and $\|u\|_{C^{0,1}(\mathbb{T}^d)}$ depends only on $1 - \eta\gamma_\varepsilon$, $\|\rho\|_{C^{0,1}}$, $\|b\|_{C^{1,1}}$, and $\|v\|_{C^{0,1}}$.*

Proof Let $\delta > 0$ and consider the viscosity regularised version of (2.14)

$$u_\delta + \gamma_\varepsilon \rho^{-2} \operatorname{div}(\rho^2 b u_\delta) - \delta \Delta u_\delta = \rho^{-1} v \quad \text{on } \mathbb{T}^d. \quad (4.7)$$

By standard elliptic PDE theory [29], (4.7) has a unique solution $u_\delta \in C^{2,\nu}(\mathbb{T}^d)$. It is a standard result in viscosity solution theory (see, e.g., [18, 10]) that $u_\delta \rightarrow u$ uniformly as $\delta \rightarrow 0$, provided $\eta\gamma_\varepsilon < 1$.

We now prove that the Lipschitz constant of u_δ is controlled independently of $\delta > 0$. The argument is standard in elliptic PDEs, and we include it for completeness. We write $c(x) = 1 + \rho^{-2} \operatorname{div}(\rho^2 b)$ for convenience, and note we have $c(x) \geq 1 - \eta\gamma_\varepsilon$. By the maximum principle, we have

$$\|u_\delta\|_{L^\infty(\mathbb{T}^d)} \leq (1 - \eta\gamma_\varepsilon)^{-1} \|\rho^{-1} v\|_{L^\infty(\mathbb{T}^d)}. \quad (4.8)$$

To bound the gradient, we differentiate both sides of (4.7) in x_i , then multiply both sides by u_{x_i} , and sum over i to obtain

$$c|\nabla u_\delta|^2 + u_\delta \nabla u_\delta \cdot \nabla c + \nabla u_\delta^T Db \nabla u_\delta + b \cdot \nabla |\nabla u_\delta|^2 - \sum_{i=1}^d u_{\delta, x_i} \Delta u_{\delta, x_i} = \nabla u_\delta \cdot \nabla (\rho^{-1} v).$$

We use the identity

$$w\Delta w = \frac{1}{2}\Delta w^2 - |\nabla w|^2 \leq \frac{1}{2}\Delta w^2$$

with $w = u_{\delta, x_i}$ to obtain

$$c|\nabla u_{\delta}|^2 + u_{\delta} \nabla u_{\delta} \cdot \nabla c + \nabla u_{\delta}^T D b \nabla u_{\delta} + b \cdot \nabla |\nabla u_{\delta}|^2 - \frac{\delta}{2} \Delta |\nabla u_{\delta}|^2 \leq \nabla u_{\delta} \cdot \nabla (\rho^{-1} v)$$

on the torus $\mathbb{T}^d$. Now, let $x_0 \in \mathbb{T}^d$ be a point where $|\nabla u_{\delta}|^2$ attains its maximum value on $\mathbb{T}^d$. Then we have $\nabla |\nabla u_{\delta}(x_0)|^2 = 0$ and $\Delta |\nabla u_{\delta}(x_0)|^2 \leq 0$, which yields

$$\begin{aligned} (1 - \eta\gamma_{\varepsilon})|\nabla u_{\delta}(x_0)|^2 &\leq c(x_0)|\nabla u_{\delta}(x_0)|^2 \\ &\leq -u_{\delta} \nabla u_{\delta} \cdot \nabla c - \nabla u_{\delta}^T D b \nabla u_{\delta} + \nabla u_{\delta} \cdot \nabla (\rho^{-1} v)|_{x_0} \\ &\leq C(\|u_{\delta}\|_{L^{\infty}(\mathbb{T}^d)} + 1)|\nabla u_{\delta}(x_0)| + \|D b\|_{L^{\infty}(\mathbb{T}^d)}|\nabla u_{\delta}(x_0)|^2, \end{aligned}$$

where C depends on ρ, c and v . Since $\|D b\|_{L^{\infty}(\mathbb{T}^d)} \leq (1 - \eta\gamma_{\varepsilon})/2$, we can apply Cauchy's inequality with ε to obtain

$$|\nabla u_{\delta}(x_0)|^2 \leq C(1 - \eta\gamma_{\varepsilon})^{-2}(\|u_{\delta}\|_{L^{\infty}(\mathbb{T}^d)} + 1)^2,$$

which completes the proof. $\square$

We now give the proof of Theorem 2.5. The proof follows the method of doubling the variables, which is used in viscosity solution theory for proving the comparison principle and establishing error estimates [18, 10]. For the reader's convenience, we review the definition of viscosity solution in Section .

Proof of Theorem 2.5 First, by Lemma 4.1 we have that u_n is uniformly bounded, independent of n , with probability at least $1 - Cn \exp(-c n h^{d+2}(1 - \eta\gamma_{\varepsilon})^2)$ provided $\varepsilon + h$ is sufficiently small. We assume these conditions throughout the rest of the proof.

Define the doubling-of-variables function

$$\Phi(x, y) := u_n(x) - u(y) - \frac{\theta}{2}|x - y|^2 \text{ on } X_n \times \mathbb{T}^d,$$

which has a maximum at $(x_n, y_n) \in X_n \times \mathbb{T}^d$. Since $\Phi(x_n, y_n) \geq \Phi(x_n, x_n)$, we have

$$u_n(x_n) - u(y_n) - \frac{\theta}{2}|x_n - y_n|^2 \geq u_n(x_n) - u(x_n).$$

By Lemma 4.3, u is Lipschitz and so

$$\frac{\theta}{2}|x_n - y_n|^2 \leq u(x_n) - u(y_n) \leq C|x_n - y_n|.$$

Hence, we have the bound $|x_n - y_n| \leq \frac{C}{\theta}$.

Define $\psi(x) = \frac{\theta}{2}|x - y_n|^2$ and $\xi_n = u_n(x_n) - \frac{\theta}{2}|x_n - y_n|^2$. By the definition of (x_n, y_n) , $u_n - \psi$ has a maximum at x_n relative to X_n . Since we have $u_n(x_n) = \psi(x_n) + \xi_n$ we see that $u_n \leq \psi + \xi_n$ on X_n and so by Proposition 4.2 we have $\mathcal{L}_n u_n(x_n) \leq \mathcal{L}_n(\psi + \xi_n)(x_n)$. It follows that

$$\begin{aligned}\frac{nh^d}{d_{x_n}}v(x_n) &= u_n(x_n) - \gamma \mathcal{L}_n u_n(x_n) \\ &\geq u_n(x_n) - \gamma \mathcal{L}_n(\psi + \xi_n)(x_n).\end{aligned}$$

By Theorem 3.5, Lemma 3.3 we have

$$\rho(x_n)^{-1}v(x_n) \geq u_n(x_n) + \gamma_\varepsilon(\rho^{-2}\operatorname{div}(\rho^2 b)u_n + \nabla\psi \cdot b)|_{x_n} - C\theta(\delta + \varepsilon + \gamma_h), \quad (4.9)$$

with probability at least $1 - Cn \exp(-cnh^{d+2}\delta^2)$.

We now define $\varphi(y) = -\frac{\theta}{2}|x_n - y|^2$. Then $u - \varphi$ attains its minimum over $\mathbb{T}^d$ at y_n , and so by the definition of viscosity supersolution we have

$$\rho(y_n)^{-1}v(y_n) \leq u(y_n) + \gamma_\varepsilon(\rho^{-2}\operatorname{div}(\rho^2 b)u + \nabla\varphi \cdot b)|_{y_n}. \quad (4.10)$$

Write $c(x) = 1 + \gamma_\varepsilon \rho^{-1}\operatorname{div}(\rho^2 b)$ and note that $\nabla\psi(x_n) = \nabla\varphi(y_n) = \theta(x_n - y_n)$ is bounded. Now subtract (4.9) from (4.10) to find that

$$c(x_n)u_n(x_n) - c(y_n)u(y_n) \leq C|x_n - y_n| + C\theta(\delta + \varepsilon + \gamma_h),$$

where $C > 0$ depends on the Lipschitz constants of ρ , v and b , and the positive lower bound for ρ . We now write

$$c(x_n)u_n(x_n) - c(y_n)u(y_n) = c(x_n)(u_n(x_n) - u(y_n)) + (c(x_n) - c(y_n))u(y_n)$$

to obtain

$$c(x_n)(u_n(x_n) - u(y_n)) \leq C|x_n - y_n| + C\theta(\delta + \varepsilon + \gamma_h),$$

where C depends additionally on $\|u\|_\infty$. Since $c(x_n) \geq 1 - \eta\gamma_\varepsilon > 0$ and $|x_n - y_n| \leq C/\theta$ we have

$$\max_{x \in \mathcal{X}_n}(u_n(x) - u(x)) \leq u_n(x_n) - u(y_n) \leq \frac{C}{\theta} + C\theta(\delta + \varepsilon + \gamma_h).$$

Optimising over $\theta > 0$ yields

$$\max_{x \in \mathcal{X}_n}(u_n(x) - u(x)) \leq C\sqrt{\delta + \varepsilon + \gamma_h}.$$

The proof of the bound in the other direction on $u - u_n$ is analogous, except we use the auxiliary function

$$\Phi(x, y) = u(x) - u_n(y) - \frac{\theta}{2}|x - y|^2. \quad \square$$

5 Numerical experiments

We now turn to some brief numerical experiments to illustrate our main results. We present numerical experiments confirming the convergence rates and parameter scalings from Theorems 2.3 and 2.5 in Sections 5.1 and 5.2, and we give an application of PageRank to data depth and high dimensional medians. The code for the numerical experiments was implemented in Python 3.7.9 and is available on GitHub: <https://github.com/jwcalder/ContinuumPageRank>.

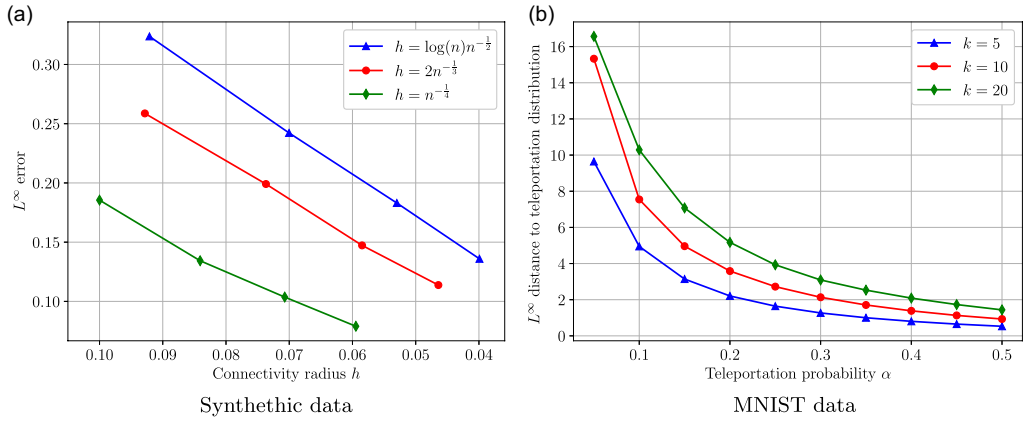


FIGURE 2. Left: For an explicit solution to the continuum limit PDE (2.10), we plot the connectivity radius h vs. the L^∞ norm of the solution error as the number of nodes increases for the three different dependencies $h = h(n)$ noted in the legend. The experiment demonstrates the linear convergence rate in our main result Theorem 2.3. Right: A computational experiment on MNIST data illustrating, for different numbers of neighbours k , how the L^∞ norm to the uniform teleportation distribution depends on the transportation probability α .

5.1 Convergence rates and parameter scalings

To test the linear convergence rates in our main result Theorem 2.3, we work with an explicit solution of the continuum PDE (2.10). We take the two-dimensional torus $\mathbb{T}^2$ and set

$$u(x) = 2 - (\cos(2\pi x_1) + \cos(2\pi x_2)),$$

and

$$v(x) = 2 - \left(1 + \frac{1}{2}\gamma_h\pi^2\right)(\cos(2\pi x_1) + \cos(2\pi x_2)),$$

where γ_h is given in (2.11). This pair of functions satisfies the PDE (2.10) with $b = 0$, $\sigma_\Phi = \frac{1}{4}$ and $\rho = 1$. The corresponding discrete PageRank problem is to take the points X_n to be an *i.i.d.* sample of size n uniformly distributed on the torus $\mathbb{T}^2$ and construct a geometric graph over the data with edge weights given by (2.3) with $b(x) = 0$, $\varphi(t) = \frac{1}{\pi}$ for $0 \leq t \leq 1$, and $\Phi(t) = 0$ otherwise.

We ran experiments with $n = 10^4, 2 \times 10^4, 4 \times 10^4$ and $n = 8 \times 10^4$, and three different choices for the connectivity length scale h : $h = \log(n)n^{-\frac{1}{2}}$, $h = 2n^{-\frac{1}{3}}$ and $h = n^{-\frac{1}{4}}$. To keep γ_h roughly constant, we set $\alpha = Ch^2$ in each case, with $C = 30, 20$ and 10 , respectively. We ran 100 trials for each combination of n , h and α and computed the average L^∞ error over all trials. The experiments were conducted on a 24-core machine with 3.8 Gz processors and 256 GB of RAM. For the largest number of vertices, $n = 8 \times 10^4$, each trial took approximately 6 min on a single core and used 8 GB of RAM.

Theorem 2.3 guarantees a linear $O(h)$ convergence rate for a much larger length scale $h \gg \log(n)^{\frac{1}{6}}n^{-\frac{1}{6}}$ (due to (2.20)), but the graphs are too dense at this length scale to perform experiments with large n . Nevertheless, we observe linear $O(h)$ convergence rates at all three smaller length scales, as is shown in the plot in Figure 2(a). It would be an interesting problem for future work to prove the linear rate from Theorem 2.3 at these smaller length scales.

To test how the parameter scalings suggested by our main results transfer to real data, we ran an experiment with the MNIST dataset [39] to compare how the distance between the PageRank vector and the teleportation distribution depends on α . The MNIST dataset contains 70,000 28×28 pixel greyscale images of handwritten digits 0–9 (see Figure 3 for an example of some MNIST images). We used all 70,000 images to construct the graph and flattened the images into vectors in $\mathbb{R}^{784}$. We used a k -nearest neighbour graph with edge weight between images i, j is given by

$$w_{i,j} = \exp\left(\frac{-4|x_i - x_j|^2}{d_k(x_i)^2}\right),$$

where $d_k(x_i)$ is the Euclidean distance between image x_i and its k th nearest neighbour. A k -nearest neighbour graph is naturally directed (i.e., asymmetric), since the k -nearest neighbour relation is not symmetric. We use the uniform teleportation probability distribution v .

The continuum PDE (2.10) suggests that the difference between the PageRank vector u and the teleportation distribution v is proportional to $\gamma_n = (\alpha^{-1} - 1)h^2$. In Figure 2(b), we show the L^∞ distance between the teleportation distribution and PageRank vector as a function of α for k -nearest neighbour graphs with $k = 10, 20, 30$. The relationship between the error and α is well-approximated by the predicted scaling of α^{-1} (a regression yields a power law α^{-p} with $p = 1.21, 1.27$ and $p = 1.05$ for $k = 10, 20$ and $k = 30$, respectively). We also see that the error is increasing with k , as expected, since h grows with k . In practice, it is important to choose α small enough so that the PageRank vector does not simply copy the teleportation distribution and instead explores the geometry of the graph, while not choosing α so small that the PageRank iterations are slow to converge.

We also mention that we ran the same experiment on the Fashion MNIST dataset [53], and obtained very similar results. The Fashion MNIST dataset is a drop-in replacement for MNIST, where the classes are types of clothing, and each data point is a 28×28 pixel greyscale image of a clothing item from a fashion catalog (see Figure 4 for an example of some Fashion MNIST images).

5.2 PageRank for data depth

In this section, we present an application of using PageRank for computing a notion of data depth and high dimensional medians. We ran experiments on the MNIST and Fashion MNIST datasets, using the same graph construction as described in Section 5.1, with $k = 10$ Euclidean nearest neighbours. We used the localised PageRank algorithm, focused on each class separately, with $\alpha = 0.05$. That is, we ran one experiment for each class with the teleportation distribution taken as the characteristic function of that class. This generates a ranking of all images relative to each class, which can be interpreted as a notion of in-class data depth, with the highest ranked images corresponding to a generalisation of median images.

In Figures 3 and 4, we show the top 11 ranked images for each class (on the right), alongside a random selection of images (on the left). We note that the highest ranked images are the most ‘typical’ for each class, which can be interpreted as median images. There is very little variation amongst the highest ranked digits, as compared to random digits. The same is observed for Fashion MNIST; the ranked images display less variation, have fewer patterns and are more solid-coloured.

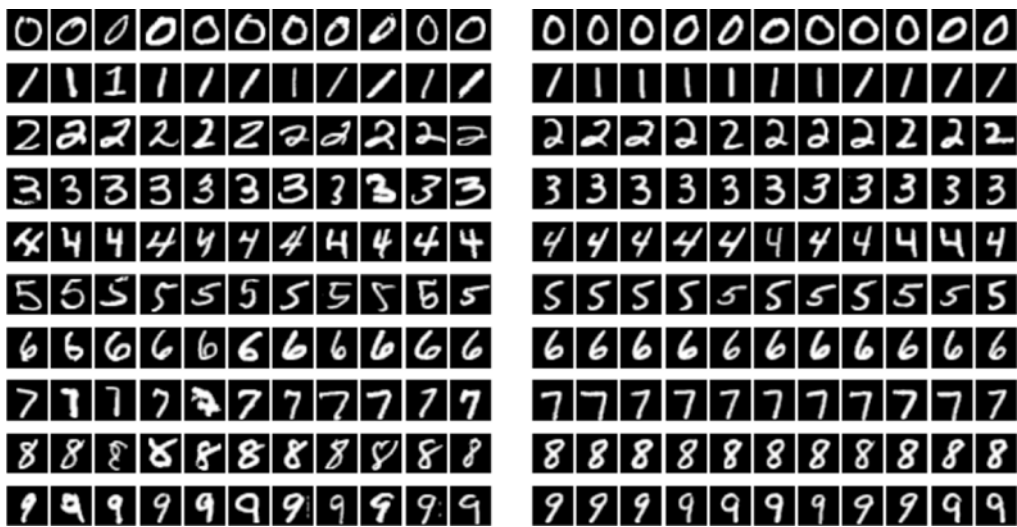


FIGURE 3. Left: A selection of 11 random samples per each digit from the MNIST dataset. Right: The highest ranked images per digit, using localised PageRank with $\alpha = 0.05$.

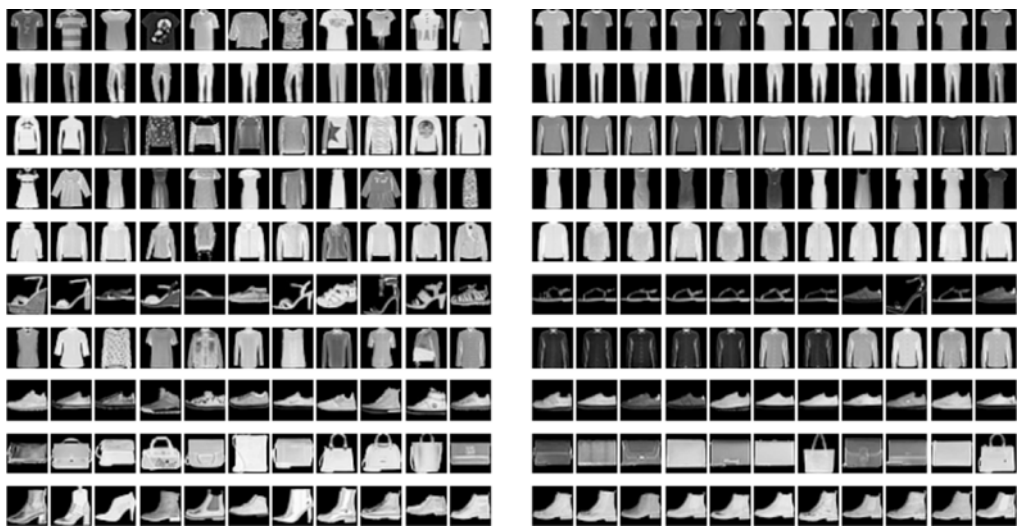


FIGURE 4. Left: A selection of 11 random images per class from the Fashion MNIST dataset, separated by type (i.e., handbag, shirt, pants, etc.). Right: The highest ranked images per type of item, using localised PageRank with $\alpha = 0.05$. Notice that the highest ranked images are more pattern-free when compared to the random images.

6 Conclusion

In this paper, we established a new framework for rigorously studying continuum limits for discrete learning problems on directed graphs. We introduced the *random directed geometric graph* and showed how to prove a continuum limit with this graph model for the PageRank

algorithm. Depending on parameter scalings, the continuum limit is either a second-order reaction–diffusion–advection equation or a first-order reaction–advection equation, whose solution can be interpreted in the viscosity sense. We also presented some numerical experiments confirming our theoretical results and exploring applications of the PageRank algorithm to multi-variate data depth.

Acknowledgements

A. Yuan and J. Calder were partially supported by NSF-DMS grants 1713691 and 1944925, and a University of Minnesota Grant in Aid award. J. Calder was also partially supported by the Alfred P. Sloan Foundation. B. Osting was partially supported by NSF DMS 16-19755 and DMS 17-52202.

Conflicts of interest

None.

References

- [1] ANDO, R. K. & ZHANG, T. (2007) Learning on graph with Laplacian regularization. In: *Advances in Neural Information Processing Systems*, Vol. 19, p. 25.
- [2] BELKIN, M. & NIYOGI, P. (2002) Laplacian eigenmaps and spectral techniques for embedding and clustering. In: *Advances in Neural Information Processing Systems*, pp. 585–591.
- [3] BELKIN, M. & NIYOGI, P. (2003) Using manifold structure for partially labeled classification. In: *Advances in Neural Information Processing Systems*, pp. 953–960.
- [4] BELKIN, M. & NIYOGI, P. (2005) Towards a theoretical foundation for Laplacian-based manifold methods. In: *International Conference on Computational Learning Theory*, Springer, pp. 486–500.
- [5] BELKIN, M. & NIYOGI, P. (2007) Convergence of Laplacian eigenmaps. In: *Advances in Neural Information Processing Systems*, pp. 129–136.
- [6] BERTOZZI, A. L. & FLENNER, A. (2012) Diffuse interface models on graphs for classification of high dimensional data. *Multiscale Model. Simul.* **10**(3), 1090–1118.
- [7] BOUSQUET, O., CHAPPELLE, O. & HEIN, M. (2004) Measure based regularization. In: *Advances in Neural Information Processing Systems*, pp. 1221–1228.
- [8] BURAGO, D., IVANOV, S. & KURYLEV, Y. (2014) A graph discretization of the Laplace–Beltrami operator. *J. Spectr. Theory* **4**(4), 675–714.
- [9] CALDER, J. (2018) The game theoretic p-Laplacian and semi-supervised learning with few labels. *Nonlinearity* **32**(1), 301.
- [10] CALDER, J. (2018) Lecture notes on viscosity solutions. Online Lecture Notes: http://www-users.math.umn.edu/~jwcalder/viscosity_solutions.pdf.
- [11] CALDER, J. (2019) Consistency of Lipschitz learning with infinite unlabeled data and finite labeled data. *SIAM J. Math. Data Sci.* **1**(4), 780–812.
- [12] CALDER, J., ESEDOGLU, S. & HERO, A. O. A Hamilton–Jacobi equation for the continuum limit of nondominated sorting. *SIAM J. Math. Anal.* **46**(1), 603–638 (2014).
- [13] CALDER, J. & GARCÍA TRILLOS, N. (2019) Improved spectral convergence rates for graph Laplacians on ε -graphs and k-NN graphs. arXiv preprint.
- [14] CALDER, J. & SLEPČEV, D. (2019) Properly-weighted graph Laplacian for semi-supervised learning. *Appl. Math. Optim. Spec. Issue Optim. Data Sci.* **82**, 1111–1159.
- [15] CALDER, J., SLEPČEV, D. & THORPE, M. (2020) Rates of convergence for Laplacian semi-supervised learning with low labeling rates. arXiv:2006.02765.

- [16] CALDER, J. & SMART, C. K. (2018) The limit shape of convex hull peeling. arXiv preprint arXiv:1805.08278.
- [17] COIFMAN, R. R. & LAFON, S. (2006) Diffusion maps. *Appl. Comput. Harmon. Anal.* **21**(1), 5–30.
- [18] CRANDALL, M. G., ISHII, H. & LIONS, P.-L. (1992) User's guide to viscosity solutions of second order partial differential equations. *Bull. Am. Math. Soc.* **27**(1), 1–67.
- [19] DUNSON, D. B., WU, H.-T. & WU, N. (2019) Diffusion based Gaussian process regression via heat kernel reconstruction. arXiv preprint arXiv:1912.05680.
- [20] EL ALAOUI, A., CHENG, X., RAMDAS, A., WAINWRIGHT, M. J. & JORDAN, M. I. (2016) Asymptotic behavior of ℓ_p -based Laplacian regularization in semi-supervised learning. In: *Conference on Learning Theory*, pp. 879–906.
- [21] EVANS, L. (2010) *Partial Differential Equations*. Graduate Studies in Mathematics, American Mathematical Society.
- [22] FLORES, M., CALDER, J. & LERMAN, G. (2018) Algorithms for ℓ_p -based semi-supervised learning on graphs. arXiv preprint.
- [23] GARCIA-CARDONA, C., MERKURJEV, E., BERTOZZI, A. L., FLENNER, A. & PERCUS, A. G. (2014) Multiclass data segmentation using diffuse interface methods on graphs. *IEEE Trans. Pattern Anal. Mach. Intell.* **36**(8), 1600–1613.
- [24] GARCÍA TRILLOS, N. (2019) Variational limits of k-NN graph-based functionals on data clouds. *SIAM J. Math. Data Sci.* **1**(1), 93–120.
- [25] GARCÍA TRILLOS, N., GERLACH, M., HEIN, M. & SLEPČEV, D. (2020) Error estimates for spectral convergence of the graph Laplacian on random geometric graphs toward the Laplace–Beltrami operator. *Found. Comput. Math.* **20**(4), 827–887.
- [26] GARCÍA TRILLOS, N. & MURRAY, R. W. (2020) A maximum principle argument for the uniform convergence of graph Laplacian regressors. *SIAM J. Math. Data Sci.* **2**(3), 705–739.
- [27] GARCÍA TRILLOS, N. & SLEPČEV, D. (2018) A variational approach to the consistency of spectral clustering. *Appl. Comput. Harm. Anal.* **45**(2), 239–281.
- [28] GARCÍA TRILLOS, N., SLEPČEV, D., VON BRECHT, J., LAURENT, T. & BRESSON, X. (2016) Consistency of Cheeger and ratio graph cuts. *J. Mach. Learn. Res.* **17**(1), 6268–6313.
- [29] GILBARG, D. & TRUDINGER, N. (2001) *Elliptic Partial Differential Equations of Second Order*. Classics in Mathematics, U.S. Government Printing Office.
- [30] GLEICH, D. F. (2015) PageRank beyond the web. *SIAM Rev.* **57**(3), 321–363.
- [31] HAVELIWALA, T. & KAMVAR, S. (2003) *The Second Eigenvalue of the Google Matrix*. Technical report, Stanford.
- [32] HE, J., LI, M., ZHANG, H.-J., TONG, H. & ZHANG, C. (2004) Manifold-ranking based image retrieval. In: *Proceedings of the 12th Annual ACM International Conference on Multimedia*, ACM, pp. 9–16.
- [33] HE, J., LI, M., ZHANG, H.-J., TONG, H. & ZHANG, C. (2006) Generalized manifold-ranking-based image retrieval. *IEEE Trans. Image Process.* **15**(10), 3170–3177.
- [34] HEIN, M., AUDIBERT, J.-Y. & VON LUXBURG, U. (2007) Graph Laplacians and their convergence on random neighborhood graphs. *J. Mach. Learn. Res.* **8**, 1325–1368.
- [35] HEIN, M., AUDIBERT, J.-Y. & VON LUXBURG, U. (2005) From graphs to manifolds—weak and strong pointwise consistency of graph Laplacians. In: *International Conference on Computational Learning Theory*, Springer, pp. 470–485.
- [36] HOFFMANN, F., HOSSEINI, B., OBERAI, A. A. & STUART, A. M. (2019) Spectral analysis of weighted Laplacians arising in data clustering. arXiv preprint arXiv:1909.06389.
- [37] LAFON, S. S. (2004) *Diffusion Maps and Geometric Harmonics*. PhD thesis, Yale University PhD dissertation.
- [38] LANGVILLE, A. N. & MEYER, C. D. (2004) Deeper inside PageRank. *Internet Math.* **1**(3), 335–380.
- [39] LECUN, Y., BOTTOU, L., BENGIO, Y. & HAFFNER, P. (1998) Gradient-based learning applied to document recognition. *Proc. IEEE* **86**(11), 2278–2324.
- [40] NG, A. Y., JORDAN, M. I. & WEISS, Y. (2002) On spectral clustering: analysis and an algorithm. In: *Advances in Neural Information Processing Systems*, pp. 849–856.

- [41] OSTING, B. & REEB, T. H. (2017) Consistency of Dirichlet partitions. *SIAM J. Math. Anal.* **49**(5), 4251–4274.
- [42] SHI, J. & MALIK, J. (2000) Normalized cuts and image segmentation. Departmental Papers (CIS), p. 107.
- [43] SHI, Z. (2015) Convergence of Laplacian spectra from random samples. arXiv preprint arXiv:1507.00151.
- [44] SHI, Z., WANG, B. & OSHER, S. J. (2018) Error estimation of weighted nonlocal Laplacian on random point cloud. arXiv preprint arXiv:1809.08622.
- [45] SHNITZER, T., BEN-CHEN, M., GUIBAS, L., TALMON, R. & WU, H.-T. (2019) Recovering hidden components in multimodal data with composite diffusion operators. *SIAM J. Math. Data Sci.* **1**(3), 588–616.
- [46] SINGER, A. (2006) From graph to manifold Laplacian: the convergence rate. *Appl. Comput. Harmon. Anal.* **21**(1), 128–134.
- [47] SLEPČEV, D. & THORPE, M. (2019) Analysis of p -Laplacian regularization in semi-supervised learning. *SIAM J. Math. Anal.* **51**(3), 2085–2120.
- [48] SZUMMER, M. & JAAKKOLA, T. (2002) Partially labeled classification with Markov random walks. In: *Advances in Neural Information Processing Systems*, pp. 945–952.
- [49] TING, D., HUANG, L. & JORDAN, M. (2010) An analysis of the convergence of graph Laplacians. In: *International Conference on Machine Learning (ICML)*.
- [50] TRILLOS, N. G. & SLEPČEV, D. (2016) Continuum limit of total variation on point clouds. *Arch. Ration. Mech. Anal.* **220**(1), 193–241.
- [51] VON LUXBURG, U., BELKIN, M. & BOUSQUET, O. (2008) Consistency of spectral clustering. *Ann. Stat.* **36**(2), 555–586.
- [52] WANG, Y., CHEEMA, M. A., LIN, X. & ZHANG, Q. (2013) Multi-manifold ranking: using multiple features for better image retrieval. In: *Pacific-Asia Conference on Knowledge Discovery and Data Mining*, Springer, pp. 449–460.
- [53] XIAO, H., RASUL, K. & VOLLGRAF, R. (2017) Fashion-mnist: a novel image dataset for benchmarking machine learning algorithms. arXiv:1708.07747.
- [54] XU, B., BU, J., CHEN, C., CAI, D., HE, X., LIU, W. & LUO, J. (2011) Efficient manifold ranking for image retrieval. In: *Proceedings of the 34th International ACM SIGIR Conference on Research and Development in Information Retrieval*, ACM, pp. 525–534.
- [55] YANG, C., ZHANG, L., LU, H., RUAN, X. & YANG, M.-H. (2013) Saliency detection via graph-based manifold ranking. In: *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition*, pp. 3166–3173.
- [56] ZHOU, D., BOUSQUET, O., LAL, T. N., WESTON, J. & SCHÖLKOPF, B. (2004) Learning with local and global consistency. *Adv. Neural Inf. Process. Syst.* **16**(16), 321–328.
- [57] ZHOU, D., HOFMANN, T. & SCHÖLKOPF, B. (2005) Semi-supervised learning on directed graphs. In: *Advances in Neural Information Processing Systems*, pp. 1633–1640.
- [58] ZHOU, D., HUANG, J. & SCHÖLKOPF, B. (2005) Learning from labeled and unlabeled data on a directed graph. In: *Proceedings of the 22nd International Conference on Machine Learning*, ACM, pp. 1036–1043.
- [59] ZHOU, D., WESTON, J., GRETTON, A., BOUSQUET, O. & SCHÖLKOPF, B. (2004) Ranking on data manifolds. In: *Advances in Neural Information Processing Systems*, Vol. 16, pp. 169–176.
- [60] ZHOU, X., BELKIN, M. & SREBRO, N. (2011) An iterated graph Laplacian approach for ranking on manifolds. In: *Proceedings of the 17th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, ACM, pp. 877–885.
- [61] ZHU, X., GHAHRAMANI, Z. & LAFFERTY, J. (2003) Semi-supervised learning using Gaussian fields and harmonic functions. In: *International Conference on Machine Learning*, Vol. 3, pp. 912–919.
- [62] ZOSSO, D. & OSTING, B. (2016) A minimal surface criterion for graph partitioning. *Inverse Probl. Imaging* **10**(4), 1149–1180.

Appendix A A more general operator

In this section, we consider the more general case where the matrix B in the definition of the weights for a random directed geometric graph (see Section 2.1 for definitions) is not the identity matrix. This case turns out to be difficult to interpret, and so we omit the details of extending our main results to this setting, but include the discussion below for completeness.

Assume that $B(x) \in \mathbb{R}^{d \times d}$ has bounded, Lipschitz continuous entries and has full rank for all $x \in \mathbb{T}^d$. Then, following the arguments of the previous subsection, the continuum limit operator is

$$L^B u(x) := -\gamma_h \left(\sigma_\Phi \operatorname{Tr} (B(x)^{-1} B(x)^{-1} \nabla^2 u(x)) + \sum_{i,j} u_{x_i x_j}(x) \int \Phi(|z|) f(x, z)_i f(x, z)_j dz \right) + \sum_i b^i(x) u_{x_i}(x) + c(x) u(x),$$

where the first-order term is

$$\begin{aligned} & \sum_i b^i(x) u_{x_i}(x) \\ &= \gamma_\varepsilon \nabla u(x) \cdot \nabla b(x) - \gamma_h \left(\sigma_\Phi (B(x)^{-1})^2 \nabla \log \rho(x) - F \right) \cdot \nabla u(x) \\ & - \gamma_h \int \Phi(|z|) (\nabla u(x) \cdot f(x, z)) (\nabla \log \rho(x) \cdot f(x, z)) dz \\ & + 2\gamma_h \sum_i u_{x_i}(x) \int \Phi(|z|) ((B(x)^{-1} z)_i + f(x, z)_i) \operatorname{Tr} (B(x)^{-1} D B(x) B(x)^{-1} z) dz, \end{aligned}$$

and the zeroth order term is

$$\begin{aligned} c(x) u(x) = & \left(1 + \gamma_\varepsilon (\rho(x)^{-2} \operatorname{div}(\rho(x)^2 b(x)) - \operatorname{Tr} (B(x)^{-1} D B(x) b(x))) \right. \\ & \left. + \gamma_h (\nabla \log \rho(x) \cdot F - H_\Phi) \right) u(x). \end{aligned}$$

We use $f(x, z)$ to represent the vector

$$(B(x)^{-1} D B(x) B(x)^{-1} z) x$$

and F to represent the vector

$$B(x)^{-1} D B(x) B(x)^{-1} z B(x)^{-1} D B(x) x B(x)^{-1} z - (B(x)^{-1} D B(x) B(x)^{-1} z)^2 x.$$

To clarify the notation, we use $D B(x)$ to represent the tensor where the ij th term is $\nabla B^{ij}(x)$. For instance, the ij th component of the term $D B(x) B(x)^{-1} z$ is the dot product of the gradient of the ij th component of B with the vector $B(x)^{-1} z$, i.e.,

$$\nabla B^{ij}(x) \cdot B(x)^{-1} z.$$

In the zeroth order term, we use H_Φ to represent the integral

$$\begin{aligned} \int \Phi(|z|) & \left(2\text{Tr}^2 \left(B(x)^{-1} DB(x) B(x)^{-1} z \right) - \text{Tr} \left(\left(B(x)^{-1} DB(x) B(x)^{-1} z \right)^2 \right) \right. \\ & - \frac{1}{2} \text{Tr} \left(B(x)^{-1} (B(x)^{-1} z D^2 B(x) B(x)^{-1} z) \right) + B(x)^{-1} DB(x) B(x)^{-1} z B(x)^{-1} z \\ & + B(x)^{-1} D^2 B(x) z B(x)^{-1} z - B(x)^{-1} DB(x) B(x)^{-1} DB(x) B(x)^{-1} z B(x)^{-1} z \\ & \left. - 2 \nabla \log \rho(x) \cdot f(x, z) \text{Tr} \left(B(x)^{-1} DB(x) B(x)^{-1} z \right) \right) dz. \end{aligned}$$

Due to the presence of $DB(x)$, the terms in F, f and H_Φ will vanish when B is a constant matrix, such as the identity.

Appendix B Definition of viscosity solution

Viscosity solutions are a notion of weak solution for PDEs based on the maximum principle. Viscosity solutions enjoy very strong stability and uniqueness properties, and the theory is especially useful for passing from discrete to continuum limits (see, e.g., [10, 16, 12]). We review here the basic definitions of viscosity solutions of the first-order equation

$$H(x, u, \nabla u) = 0 \quad \text{in } \Omega, \quad (.1)$$

where $\Omega \subset \mathbb{R}^d$ is open. For viscosity solution on the torus $\mathbb{T}^d$, we take $\Omega = \mathbb{R}^n$ and treat functions on $\mathbb{T}^d$ as $\mathbb{Z}^d$ -periodic functions on $\mathbb{R}^d$ for defining viscosity solutions.

Definition .1 We say $u \in C(\bar{\Omega})$ is a viscosity subsolution of (.1) if for all $x \in \Omega$ and every $\varphi \in C^\infty(\mathbb{R}^n)$ such that $u - \varphi$ has a local maximum at x we have

$$H(x, u(x), \nabla \varphi(x)) \leq 0.$$

Likewise, we say that $u \in C(\bar{U})$ is a viscosity supersolution of (.1) if for all $x \in \Omega$ and every $\varphi \in C^\infty(\mathbb{R}^n)$ such that $u - \varphi$ has a local minimum at x we have

$$H(x, u(x), \nabla \varphi(x)) \geq 0.$$

Finally, we say that u is a viscosity solution of (.1) if u is both a viscosity sub- and supersolution.