

Analysis of a Two-Layer Neural Network via Displacement Convexity

Adel Javanmard*, Marco Mondelli† and Andrea Montanari‡

August 20, 2019

Abstract

Fitting a function by using linear combinations of a large number N of ‘simple’ components is one of the most fruitful ideas in statistical learning. This idea lies at the core of a variety of methods, from two-layer neural networks to kernel regression, to boosting. In general, the resulting risk minimization problem is non-convex and is solved by gradient descent or its variants. Unfortunately, little is known about global convergence properties of these approaches.

Here we consider the problem of learning a concave function f on a compact convex domain $\Omega \subset \mathbb{R}^d$, using linear combinations of ‘bump-like’ components (neurons). The parameters to be fitted are the centers of N bumps, and the resulting empirical risk minimization problem is highly non-convex. We prove that, in the limit in which the number of neurons diverges, the evolution of gradient descent converges to a Wasserstein gradient flow in the space of probability distributions over Ω . Further, when the bump width δ tends to 0, this gradient flow has a limit which is a viscous porous medium equation. Remarkably, the cost function optimized by this gradient flow exhibits a special property known as *displacement convexity*, which implies exponential convergence rates for $N \rightarrow \infty$, $\delta \rightarrow 0$.

Surprisingly, this asymptotic theory appears to capture well the behavior for moderate values of δ, N . Explaining this phenomenon, and understanding the dependence on δ, N in a quantitative manner remains an outstanding challenge.

1 Introduction

In supervised learning, we are given data $\{(y_j, \mathbf{x}_j)\}_{j \leq n}$ which are often assumed to be independent and identically distributed from a common law $\mathbb{P}$ on $\mathbb{R} \times \mathbb{R}^d$ (here $\mathbf{x}_j \in \mathbb{R}^d$ is a feature vector, and $y_j \in \mathbb{R}$ is a label or response variable). We would like to find a function $\hat{f} : \mathbb{R}^d \rightarrow \mathbb{R}$ to predict the labels at new points $\mathbf{x} \in \mathbb{R}^d$. Throughout this paper, we will quantify the quality of our prediction by square loss, hence we are interested in minimizing $R(\hat{f}) = \mathbb{E}\{(y - \hat{f}(\mathbf{x}))^2\}$.

One of the most fruitful ideas in this context is to use functions that are linear combinations of

*Data Science and Operations Department, Marshall School of Business, University of Southern California

†Department of Electrical Engineering, Stanford University

‡Department of Electrical Engineering and Department of Statistics, Stanford University

simple components:

$$\hat{f}(\mathbf{x}; \mathbf{w}) = \frac{1}{N} \sum_{i=1}^N a_i \sigma(\mathbf{x}; \mathbf{w}_i). \quad (1.1)$$

Here $\sigma : \mathbb{R}^d \times \mathbb{R}^D \rightarrow \mathbb{R}$ is a component function (a ‘neuron’ or ‘unit’ in the neural network parlance), and $\mathbf{w} = (\mathbf{w}_1, \dots, \mathbf{w}_N) \in \mathbb{R}^{D \times N}$, $\mathbf{a} = (a_1, \dots, a_N) \in \mathbb{R}^N$ are parameters to be learnt from data. Standard choices for the activation function are $\sigma(\mathbf{x}; \mathbf{w}) = (1 + \exp(-\langle \mathbf{w}, \mathbf{x} \rangle))^{-1}$ (sigmoid) or $\sigma(\mathbf{x}; \mathbf{w}) = \max(\langle \mathbf{w}, \mathbf{x} \rangle; 0)$ (ReLU). In this paper we will instead study a class of activation that depends on the difference $\mathbf{x} - \mathbf{w}$. The objective is to minimize the population (prediction) risk

$$R_N(\mathbf{a}, \mathbf{w}) = \mathbb{E} \left\{ \left[y - \frac{1}{N} \sum_{i=1}^N a_i \sigma(\mathbf{x}; \mathbf{w}_i) \right]^2 \right\}. \quad (1.2)$$

Special instantiations of this idea include (we provide only pointers to the immense literature on each topic):

- Two-layer neural networks [Ros62, AB09];
- Sparse deconvolution [Don92, CFG14];
- Kernel ridge regression and related random feature methods [CST00, RR08];
- Boosting [Sch03, Fri01, BY03].

Despite the impressive practical success of these methods, the risk function $R_N(\mathbf{w})$ is highly non-convex and little is known about global convergence of algorithms that try to minimize it (we refer to Section 2 for further discussion of the related literature).

Notable exceptions to the last statement are provided by random features and by boosting algorithms. In random feature methods, the parameters $\mathbf{w}_i$ are not optimized over (they are drawn i.i.d. from some common distribution), and the resulting risk function becomes convex in the weights $(a_1, \dots, a_N)$ to be learnt. While this is a fruitful idea, it gives up the degrees of freedom afforded by the $\mathbf{w}_i$ ’s.

Boosting overcomes non-convexity by fitting the components $\mathbf{w}_1, \dots, \mathbf{w}_N$ one at the time, sequentially. The underlying assumption is that the problem of minimizing $R_N(\mathbf{w})$ with respect to one of the hidden units $\mathbf{w}_i$ is tractable. However, this is generally not the case when the parameters $\mathbf{w}_i$ belong to a high-dimensional space.

The risk function (1.2) crystalizes a central conundrum in statistical learning. In a number of applications (especially at low noise), it is rarely the case that low prediction error can be achieved through a function that is linear in the raw covariates, e.g. $\hat{f}(x) = \langle \mathbf{w}, \mathbf{x} \rangle$. In a classical setting, the statistician would craft nonlinear features out of the covariates on the basis of expert knowledge. For the model of Eq. (1.1), this amounts to constructing vectors $\mathbf{w}_1, \dots, \mathbf{w}_N$. Statistical methods would then be confined to the convex task of fitting the coefficients $a_1, \dots, a_N$. This step is well understood from a statistical and computational perspective.

Modern machine learning approaches (boosting, neural networks, etc.) hold the promise of automatizing feature extraction, hence producing superior performances in a wide variety of applications. Unfortunately, we are still far from understanding in which cases optimizing over the

$\mathbf{w}_i$'s yields a significant improvement over –say– choosing them randomly. This central challenge intertwines statistical and computational aspects. It is not hard to see that varying the weights $\mathbf{w}_i$'s produces a significantly larger function class [Bac17]. The relevant question is what part of this class can be accessed using gradient descent or other practical algorithms.

The main objective of this paper is to introduce a nonparametric regression model in which these questions can be addressed rigorously. The model is interesting for at least two reasons: (i) From a theoretical point of view, global convergence can be proved in the limit of a large neurons. The proof relies on a mathematical mechanism that has not been explored in the statistics or machine learning literature before. (ii) From a practical point of view, the model is nontrivial enough to illustrate the potential advantage of fitting the features $\mathbf{w}_i$ (we demonstrate this numerically in Section 4.)

Let $\Omega \subset \mathbb{R}^d$ be a compact convex set with $\mathcal{C}^2$ boundary. We assume $\{(y_j, \mathbf{x}_j)\}_{j \geq 1}$ to be i.i.d. where $\mathbf{x}_j \sim \text{Unif}(\Omega)$ and

$$\mathbb{E}(y_j | \mathbf{x}_j) = f(\mathbf{x}_j), \quad (1.3)$$

with $f : \Omega \rightarrow \mathbb{R}$ a smooth function. We try to fit these data using a combination of bumps, namely

$$\hat{f}(\mathbf{x}; \mathbf{w}) = \frac{1}{N} \sum_{i=1}^N K^\delta(\mathbf{x} - \mathbf{w}_i), \quad (1.4)$$

where $K^\delta(\mathbf{x}) = \delta^{-d} K(\mathbf{x}/\delta)$, $K : \mathbb{R}^d \rightarrow \mathbb{R}_{\geq 0}$ is a first order kernel with compact support, and $\mathbf{w}_i \in \Omega^\delta$ for $i \leq N$. Here Ω^δ is a slightly smaller compact set, with $\Omega^\delta \rightarrow \Omega$ as $\delta \rightarrow 0$. (Note that in our setting the hidden units $\mathbf{w}_i$ and input data $\mathbf{x}_j$ have same dimensions, i.e., $d = D$.) We refer to Section 5 for a formal statement of our assumptions. From Eq. (1.2), we have

$$R_N(\mathbf{w}) = R_\# + \mathbb{E}\left\{\left[f(\mathbf{x}) - \frac{1}{N} \sum_{i=1}^N K^\delta(\mathbf{x} - \mathbf{w}_i)\right]^2\right\},$$

where $R_\# = \mathbb{E}[(y - f(\mathbf{x}))^2]$ and we use the fact that $\mathbb{E}[y - f(\mathbf{x}) | \mathbf{x}] = 0$. Since the constant $R_\#$ does not depend on parameters $\mathbf{w}$, it does not matter in optimizing $R_N(\mathbf{w})$ over $\mathbf{w}$ and henceforth we write, with a slight abuse of notation,

$$R_N(\mathbf{w}) = \mathbb{E}\left\{\left[f(\mathbf{x}) - \frac{1}{N} \sum_{i=1}^N K^\delta(\mathbf{x} - \mathbf{w}_i)\right]^2\right\}.$$

The model (1.4) is general enough to include a broad class of radial-basis function (RBF) networks which are known to be universal function approximators [PS91]. To the best of our knowledge, there is no result on the global convergence of stochastic gradient descent for learning RBF networks, and this paper establishes the first result of this type.

It is important to emphasize a few differences with respect to standard RBF networks. First of all, we do not require the kernel $K(\mathbf{x})$ to be radial, i.e. to depend uniquely on the norm $|\mathbf{x}|$. Second, we require K to have compact support. This is mainly a technical requirement that simplifies some arguments: we expect our results to be generalizable to kernels that decay rapidly enough. Finally, and most crucially, the form (1.4) does not include non-uniform weights for the N components.

A more standard formulation would posit $\hat{f}(\mathbf{x}; \mathbf{w}) = \sum_{i=1}^N a_i K^\delta(\mathbf{x} - \mathbf{w}_i)$ and learn the weights a_i from data, see Eq. (1.1). We deliberately set the weights to a fixed value because the risk function is convex in $\mathbf{a} = (a_i)_{i \leq N}$, and hence fitting $\mathbf{a}$'s to global optimality is ‘easy.’ Indeed, universal approximation could be achieved by keeping the centers $\mathbf{w}_i$ fixed (and sufficiently dense in Ω) and only adjusting $\mathbf{a}$. As discussed above, our focus is on the role of the $\mathbf{w}_i$'s.

Our main result is a proof that, for sufficiently large N and small δ , gradient descent algorithms converge to weights $\mathbf{w}$ with nearly optimum prediction error, provided f is strongly concave. Let us emphasize that the resulting population risk $R_N(\mathbf{w})$ is non-convex regardless of the concavity properties of f . Our proof unveils a novel mechanism by which global convergence takes place. Convergence results for non-convex empirical risk minimization are generally proved by carefully ruling out local minima in the cost function (see Section 2 for pointers to this literature). Instead we prove that, as $N \rightarrow \infty$, $\delta \rightarrow 0$, the gradient descent dynamics converges to a gradient flow in Wasserstein space, and that the corresponding cost function is ‘displacement convex.’ Breakthrough results in optimal transport theory guarantee dimension-free convergence rates for this limiting dynamics [CJM⁺01, CMV03, CMV06]. In particular, we expect the cost function $R_N(\mathbf{w})$ to have many local minima, which are however completely neglected by the gradient descent dynamics.

More specifically, our first step is to show that – for large N – the evolution of the weights $\mathbf{w}_1, \dots, \mathbf{w}_N$ under gradient descent can be replaced by the evolution of a probability distribution¹ $\rho^\delta \in \mathcal{P}_2(\Omega)$, which approximates their empirical distribution. Namely, if $(\mathbf{w}_1^k, \dots, \mathbf{w}_N^k)$ denote the weights after k iterations with step size ε , and $\hat{\rho}_k^{(N)} = \sum_{i=1}^N \delta_{\mathbf{w}_i^k} / N$ is their empirical distribution, then we have

$$\lim_{N \rightarrow \infty, \varepsilon \rightarrow 0} \hat{\rho}_{t/\varepsilon}^{(N)} = \rho_t^\delta, \quad (1.5)$$

where the limit holds in the sense of weak convergence or in W_1 distance (the two are equivalent since Ω is compact). The limit evolution $(\rho_t^\delta)_{t \geq 0}$ satisfies a partial differential equation (PDE) that can also be described as the Wasserstein W_2 gradient flow (i.e. gradient flow in $\mathcal{P}_2(\Omega)$), for the following effective risk

$$R^\delta(\rho) = \nu_0 \int_{\Omega} [f(\mathbf{x}) - K^\delta * \rho(\mathbf{x})]^2 d\mathbf{x}, \quad (1.6)$$

where $\nu_0 = 1/|\Omega|$ and $|\Omega|$ denotes the volume of the set Ω . Here $*$ denotes the usual convolution. Let us emphasize that the convergence to Wasserstein gradient flow holds regardless of the concavity of f .

The use of W_2 gradient flows to analyze two-layer neural networks was recently developed in several papers [MMN18, RVE18, CB18, SS18]. However, we cannot rely on earlier results because of the specific boundary conditions in our problem. We constrain the $\mathbf{w}_i \in \Omega^\delta$ by running projected stochastic gradient descent (SGD): at each step $\mathbf{w}_i$ moves in the direction of a stochastic gradient of $R_N(\mathbf{w})$ and then projected back to Ω^δ . This results in a PDE with Neumann boundary condition on Ω^δ , which is not covered by previous theory. We establish a quantitative version of the limit (1.5) via propagation-of-chaos techniques.

Even if the cost (1.6) is quadratic and convex in ρ , its W_2 gradient flow can have multiple fixed points, and hence global convergence cannot be guaranteed. Global convergence results were

¹Throughout, $\mathcal{P}_2(\mathcal{X})$ denotes the space of probability distributions on $\mathcal{X}$, endowed with Wasserstein metric W_2 .

proven in [MMN18] and in [CB18] by showing that, for all $t \geq 0$ ρ_t^δ has a density that is either smooth, or strictly positive everywhere. However, these convergence results are non-quantitative, and do not provide convergence rates².

Indeed, the mathematical property that controls global convergence of W_2 gradient flow is not ordinary convexity but *displacement convexity*. Roughly speaking, displacement convexity is convexity along geodesics of the W_2 metric, see Section 3.5. The risk function (1.6) is not displacement convex. Indeed, its quadratic term reads $\nu_0 \int K_\delta * K_\delta(\mathbf{x} - \mathbf{x}') \rho(\mathbf{x}) \rho(\mathbf{x}') d\mathbf{x} d\mathbf{x}'$ which is not displacement convex unless $K_\delta * K_\delta$ is convex (see Lemma H.1), which cannot be in our setting. However, for small δ , we can formally approximate $K^\delta * \rho \approx \rho$, and hence hope to replace the risk function (1.6) with a simpler one

$$R(\rho) = \nu_0 \int_{\Omega} [f(\mathbf{x}) - \rho(\mathbf{x})]^2 d\mathbf{x}. \quad (1.7)$$

Most of our technical work is devoted to making rigorous this $\delta \rightarrow 0$ approximation. Namely, we prove that, as $\delta \rightarrow 0$, $\rho_t^\delta \Rightarrow \rho_t$ where ρ_t follows the W_2 gradient flow for the risk $R(\rho)$.

Remarkably, the risk function $R(\rho)$ is strongly displacement convex (provided f is strongly concave). A long line of work in PDE and optimal transport theory establishes dimension-free convergence rates for its W_2 gradient flow [CJM⁺01, CMV03, CMV06]. Namely, if f is α -strongly concave, then $R(\rho_t) \leq R(\rho_0) e^{-2\alpha t}$. By using the approximation results outlined above, we obtain global convergence for SGD. With high probability,

$$R_N(\mathbf{w}^k) \leq R_N(\mathbf{w}^0) e^{-2\alpha k\varepsilon} + \text{err}(N, d, \varepsilon, \delta), \quad (1.8)$$

where the error term err vanishes as $N \rightarrow \infty$, $\varepsilon, \delta \rightarrow 0$ in a suitable order.

This result implies that SGD converges exponentially fast to a near-global optimum with a rate that is controlled by the convexity parameter α .

Our bounds are not sharp enough to provide quantitative control on the error term $\text{err}(N, d, \varepsilon, \delta)$, especially in high dimension. Nevertheless, the convergence rate predicted by our asymptotic theory is in excellent agreement with numerical simulations, cf. Section 4. Explaining this surprising quantitative agreement is an outstanding challenge.

2 Related literature

The present work ties in several lines of research, some of which were already mentioned in the introduction. A substantial amount of work has been devoted to analyzing two-layer neural networks and developing algorithms with convergence guarantees, see e.g. [ZSJ⁺17, Tia17, BJW18]. However these approaches are typically based on tensor factorization or similar initialization steps that are not used in practice, and do not scale well (although polynomially) in high dimension.

The landscape of empirical risk minimization was also studied in a number of papers, see e.g. [LY17, SJJ18]. However, global convergence was only proved in the extremely overparametrized regime in which the neural network essentially behaves as kernel ridge regression [DZPS18].

²An argument indicating convergence in a time polynomial in d was put forward in [WLLM18], but for a different type of continuous flow.

Classical theory of neural networks was largely devoted to the two-layer case [AB09], although the focus was on representation and approximation questions [Cyb89, Bar93], as well as on generalization error. It was already clear in that context that a two-layer network is conveniently characterized by the empirical distribution of the hidden neurons, and that it is useful to relax this from a distribution with N atoms, to a general probability measure. This representation plays an important role, for instance, in [Bar98], and was exploited again under the label of ‘convex neural networks’ in [BRV⁺06].

Over the last year, several groups independently revisited this connection, with the objective of understanding the landscape structure of two-layer networks, and the dynamics of gradient descent methods [NS17, MMN18, RVE18, SS18, CB18, MMM19]. In particular, it was proven in [MMN18] that, under certain smoothness condition on the underlying data distribution, the gradient descent evolution is well approximated by a Wasserstein gradient flow, provided that the number of neurons exceeds the data dimensions. As mentioned above, the algorithm treated here differs from the ones analyzed in earlier work, because the weights $\mathbf{w}_i$ are constrained to lie in the convex set Ω^δ . We enforce this constraint by using projected SGD, i.e. projecting at each step the weights onto the set Ω^δ . We generalize the analysis of [MMN18], obtaining convergence to a PDE with Neumann (reflecting) boundary conditions. As in [MMN18], we build on ideas that were first developed in the context of interacting particle systems [Dob79, Szn91].

The Wasserstein gradient flow approach was used in [MMN18, CB18] to establish global convergence results. However, these results fall short of our objectives for several reasons:

- The global convergence result of [CB18] rely on certain homogeneity properties of the neurons that are lacking here. We could obtain homogeneity by adding coefficients to Eq. (1.4), i.e. considering $\hat{f}(\mathbf{x}; \mathbf{w}) = \sum_{i=1}^N a_i K^\delta(\mathbf{x} - \mathbf{w}_i)$ and minimizing the risk with respect to the coefficients a_i . As mentioned above, we refrain from introducing coefficients not to oversimplify the problem: when $N \rightarrow \infty$, it is sufficient to fit the coefficients a_i to achieve vanishing risk. Fitting the a_i ’s is a least squares problem.
- Most importantly, the techniques [MMN18, CB18] do not establish any convergence rates. This is not surprising, as those results hold under weak assumptions on the data distribution and the activation function. In particular, [MMN18, CB18, MMM19] cover general risk functions of the form (1.2) under certain smoothness and boundedness conditions on σ and on the functions $V(\mathbf{w}) = -\mathbb{E}\{f(\mathbf{x})\sigma(\mathbf{x}; \mathbf{w})\}$, $U(\mathbf{w}_1, \mathbf{w}_2) = \mathbb{E}\{\sigma(\mathbf{x}; \mathbf{w}_1)\sigma(\mathbf{x}; \mathbf{w}_2)\}$. In such a general setting [MMN18] provides examples in which the Wasserstein gradient flow has multiple fixed points, which are singular with respect to the Lebesgue measure. Global convergence is established in [MMN18, CB18] by proving that PDE solution ρ_t has a strictly positive density. However, it is difficult to imagine this condition to hold in a quantitative dimension-independent manner.

In contrast, our results are a first step towards dimension-independent convergence rate, in a more restricted setting than [MMN18, CB18, MMM19].

In summary, our results do not subsume earlier work, that assumes a more general setting, but rather establish stronger results in narrower context. Indeed, we believe that specific structural conditions must be imposed on the data distribution and activation function for the Wasserstein gradient flow approach to yield quantitative convergence rates. This paper presents one specific

set of assumptions. Although our results are not strong enough to establish non-asymptotic convergence rates, they point clearly in that direction.

3 Model and assumptions

3.1 Notations

We will use lowercase boldface for vectors, e.g. $\mathbf{x}, \mathbf{y}, \dots$, uppercase for random variables, e.g. $X, Y, \dots$, and uppercase boldface for random vectors, e.g. $\mathbf{X}, \mathbf{Y}, \dots$. The scalar product of two vectors is denoted by $\langle \mathbf{x}, \mathbf{y} \rangle = \sum_{i=1}^d x_i y_i$, and the ℓ_2 norm of a vector is denoted by $|\mathbf{x}|$. The Euclidean ball in $\mathbb{R}^d$ with center $\mathbf{x}$ and radius r is denoted by $\mathbf{B}(\mathbf{x}; r)$. Given a set $\Omega \subseteq \mathbb{R}^d$, we denote by $|\Omega|$ its volume.

We will refer to several function spaces in what follows. The most common is the space of p -th integrable functions $\mathcal{L}^p(\mathcal{X})$ on a measure space $(\mathcal{X}, \mathcal{F}, \mu)$. Given a function $f : \mathcal{X} \rightarrow \mathbb{R}$, we denote by $\|f\|_{\mathcal{L}^p(\mathcal{X})}$ its $\mathcal{L}^p$ norm, namely $\|f\|_{\mathcal{L}^p(\mathcal{X})}^p = \int_{\mathcal{X}} |f(x)|^p \mu(dx)$. For $S \subseteq \mathbb{R}^m$, $\mathcal{C}^k(S)$ denotes the space of continuous functions $f : S \rightarrow \mathbb{R}$ with continuous derivatives up to order k . In particular, $\mathcal{C}(S)$ denotes the space of continuous real-valued functions defined on S . In addition, for $T \in \mathbb{R}_+$ and a metric space $\mathcal{M}$ (with distance $d_{\mathcal{M}}$), $\mathcal{C}([0, T], \mathcal{M})$ denotes the set of continuous functions $f : [0, T] \rightarrow \mathcal{M}$, endowed with the distance between two functions $f, g \in \mathcal{C}([0, T], \mathcal{M})$ defined as $d_{\mathcal{C}([0, T], \mathcal{M})}(f, g) \equiv \sup_{t \in [0, T]} d_{\mathcal{M}}(f(t), g(t))$. For a function $f : S \rightarrow \mathbb{R}$, we let $\|f\|_{\text{Lip}} \equiv \sup_{\mathbf{x} \neq \mathbf{y} \in S} |f(\mathbf{x}) - f(\mathbf{y})|/|\mathbf{x} - \mathbf{y}|$ be the Lipschitz constant of the function f . Finally, as mentioned above, $\mathcal{P}_2(\mathcal{X})$ denotes the space of probability distributions on $\mathcal{X}$, endowed with the Wasserstein metric W_2 .

Throughout the paper, we use C to denote finite constants, which can vary from point to point. When these constants can depend on some of the problem parameters, e.g. a, b, c , we will write $C(a, b, c)$. When they are absolute numerical constants, we will emphasize this by writing C_* .

3.2 Data

As mentioned above, we are given data $(y_j, \mathbf{x}_j) \sim_{\text{i.i.d.}} \mathbb{P}$ where $\mathbf{x}_j \sim \text{Unif}(\Omega)$, with $\Omega \subset \mathbb{R}^d$ a compact convex set, and $y_j = f(\mathbf{x}_j) + \varepsilon_j$, with $f : \Omega \rightarrow \mathbb{R}_{\geq 0}$. We assume the ε_j to be i.i.d. σ^2 -subgaussian random variables with $\mathbb{E}(\varepsilon_j | \mathbf{x}_j) = 0$. We assume the function f to be concave and smooth.

Our formal assumptions on the set Ω and the function f are as follows:

- (A1) $\Omega \supseteq \mathbf{B}(\mathbf{0}; r)$, with $r > 0$, is a compact convex set with $\mathcal{C}^2$ boundary.
- (A2) $f : \Omega \rightarrow \mathbb{R}_{\geq 0}$ uniformly concave, i.e., there exists $\alpha > 0$ such that

$$\langle \mathbf{y}, \nabla^2 f(\mathbf{x}) \mathbf{y} \rangle \leq -\alpha |\mathbf{y}|^2, \quad \forall \mathbf{x} \in \Omega, \mathbf{y} \in \mathbb{R}^d, \quad (3.1)$$

where $\nabla^2 f$ denotes the Hessian of f .

- (A3) $f \in \mathcal{C}^\infty(\Omega)$, with $\|f\|_{\mathcal{L}^\infty(\Omega)}, \|\nabla f\|_{\mathcal{L}^\infty(\Omega)} \leq C_*$ for an absolute constant C_* .

Without loss of generality, we can also assume that $\int_{\Omega} f(\mathbf{x}) d\mathbf{x} = 1$. As a running example, we will use $\Omega = \mathbf{B}(\mathbf{0}; r)$, where we remind r is defined in Assumption (A1).

Remark 3.1. The assumption $\mathbf{x}_j \sim \text{Unif}(\Omega)$ is quite strong but simplifies our analysis. We believe our approach can be generalized to a broader family of probability distribution for the covariates $\mathbf{x}_j$, but defer these generalizations to future work.

3.3 Neural network and SGD

Let $K \in \mathcal{C}^2(\mathbb{R}^d)$ be a non-negative symmetric first order kernel with compact support. Formally, we assume that

$$(A4) \quad \int K(\mathbf{x}) d\mathbf{x} = 1, \quad K(\mathbf{x}) \geq 0, \quad \int K(\mathbf{x}) \mathbf{x} d\mathbf{x} = 0, \quad (3.2)$$

$$K(-\mathbf{x}) = K(\mathbf{x}), \quad \text{supp}(K) \subseteq \mathbf{B}(\mathbf{0}, c_0). \quad (3.3)$$

The assumptions of symmetry and compact support are not crucial, but simplify some of the technical details later. We will further assume $\|\nabla K\|_{\mathcal{L}^\infty(\mathbb{R}^d)}$, $\|\nabla^2 K\|_{\mathcal{L}^\infty(\mathbb{R}^d)}$ and c_0 to be independent of the ambient dimension d . Notice that this requirement follows from the differentiability and compact support assumptions if $K(\mathbf{x}) = \kappa(\|\mathbf{x}\|_2)$ is a radial function.

For $\delta > 0$, let $K^\delta(\mathbf{x}) = \delta^{-d}K(\mathbf{x}/\delta)$. We try to fit the function (1.4) with parameters $\mathbf{w} = (\mathbf{w}_1, \dots, \mathbf{w}_N)$. These parameters are constrained to $\mathbf{w}_i \in \Omega^\delta$ which is a suitable scaling of Ω , as defined in the following. Given $\delta < r/c_0$, with r defined in (A1), define

$$\Omega^\delta = \lambda_\delta \Omega,$$

where

$$\lambda_\delta = \sup \{ \lambda \geq 0 : \lambda \Omega \oplus \mathbf{B}(\mathbf{0}, c_0 \delta) \subseteq \Omega \}. \quad (3.4)$$

For two sets $A, B \subseteq \mathbb{R}^d$, their Minkowski sum is defined as $A \oplus B = \{\mathbf{x} + \mathbf{y} : \mathbf{x} \in A, \mathbf{y} \in B\}$. Note that $\lambda_\delta \in [0, 1]$ for all δ . Furthermore, $\Omega \supseteq \mathbf{B}(\mathbf{0}; r)$ implies $\lambda_\delta > 0$ for all $\delta < r/c_0$. Finally, $\lambda_{\delta=0} = 1$, whence $\Omega^{\delta=0} = \Omega$. In our running example, $\Omega^\delta = \mathbf{B}(\mathbf{0}; r - c_0 \delta)$ is a ball of slightly smaller radius. Clearly, since Ω is convex, Ω^δ is convex as well.

We use stochastic gradient descent to minimize the population risk (1.2). At each step, we use a new data point $(y_k, \mathbf{x}_k)$, thus the sample size is equal to the number of iterations of the algorithm. Assuming for simplicity constant step size $\varepsilon > 0$, we update the parameters by

$$\mathbf{w}_i^{k+1} = \mathbf{P} \left\{ \mathbf{w}_i^k - \varepsilon \nabla K^\delta(\mathbf{x}_{k+1} - \mathbf{w}_i^k) \left(y_{k+1} - \hat{f}(\mathbf{x}_{k+1}; \mathbf{w}^k) \right) + \sqrt{2\varepsilon\tau} \mathbf{g}_i^{k+1} \right\}. \quad (3.5)$$

Here $\mathbf{g}_i^{k+1} \sim \mathbf{N}(0, \mathbf{I}_d)$ is Gaussian noise which we take to be i.i.d. across time and neuron indices, k and i , and $\mathbf{P}$ is the orthogonal projector onto Ω^δ :

$$\mathbf{P}(\mathbf{z}) = \arg \min \{ |\mathbf{z} - \mathbf{x}| : \mathbf{x} \in \Omega^\delta \}. \quad (3.6)$$

The noise term $\sqrt{2\varepsilon\tau} \mathbf{g}_i^{k+1}$ is added mainly for technical reasons. Namely, it allows us to control the smoothness of the solutions of the resulting PDE. In simulations we do not find it useful, and we believe that a more careful analysis would be able to establish smoothness without the noise term.

Again, in our running example, we have

$$P(\mathbf{z}) = \begin{cases} \mathbf{z} & \text{if } |\mathbf{z}| \leq r - c_0\delta, \\ (r - c_0\delta)\mathbf{z}/|\mathbf{z}| & \text{if } |\mathbf{z}| > r - c_0\delta. \end{cases} \quad (3.7)$$

We initialize SGD with $(\mathbf{w}_i^0)_{i \leq N} \sim \text{i.i.d. } \rho_{\text{init}}^\delta \in \mathcal{P}_2(\Omega^\delta)$, where $\rho_{\text{init}}^\delta$ is a scaling of a fixed distribution $\rho_{\text{init}} \in \mathcal{P}_2(\Omega)$, i.e. $\rho_{\text{init}}^\delta(S) = \rho_{\text{init}}(S/\lambda_\delta)$. We assume that the initialization is smooth:

$$(A5) \quad \rho_{\text{init}} \in \mathcal{C}^\infty(\Omega^\delta).$$

3.4 PDE Model, $\delta > 0$

In the $N \rightarrow \infty$ limit the population risk is approximated by the effective risk $R^\delta : \mathcal{P}_2(\Omega^\delta) \rightarrow \mathbb{R}$ defined in Eq. (1.6). We emphasize that ρ is a probability distribution supported on Ω^δ . Note that

$$\inf_{\rho} R^\delta(\rho) \leq R^\delta(f) = \nu_0 \int_{\Omega} [f(\mathbf{x}) - K^\delta * f(\mathbf{x})]^2 d\mathbf{x}. \quad (3.8)$$

In particular $\lim_{\delta \rightarrow 0} \inf_{\rho \in \mathcal{P}_2(\Omega)} R^\delta(\rho) = 0$.

Our first main result is that the dynamics of SGD is well approximated by the following PDE (see Section 5.1 for a formal statement):

$$\begin{aligned} \partial_t \rho_t(\mathbf{w}) &= \nabla \cdot (\rho_t(\mathbf{w}) \nabla \Psi(\mathbf{w}; \rho_t)) + \tau \Delta \rho_t(\mathbf{w}), \\ \Psi(\mathbf{w}; \rho) &\equiv -\nu_0 K^\delta * f(\mathbf{w}) + \nu_0 K^\delta * K^\delta * \rho(\mathbf{w}), \end{aligned} \quad (3.9)$$

with initial and boundary conditions

$$\begin{aligned} \rho_0 &= \rho_{\text{init}}^\delta, \\ \langle \mathbf{n}(\mathbf{w}), \rho_t(\mathbf{w}) \nabla \Psi(\mathbf{w}; \rho_t) + \tau \nabla \rho_t(\mathbf{w}) \rangle &= 0 \quad \forall \mathbf{w} \in \partial\Omega^\delta, \end{aligned} \quad (3.10)$$

where $\mathbf{n}(\mathbf{x})$ denotes the inward normal vector to $\partial\Omega^\delta$ at $\mathbf{x}$.

A rigorous definition of solutions of this PDE, along with some of their properties, is given in Appendix B. In Appendix C, we discuss the connection between the PDE (3.9) and the so-called “nonlinear dynamics”, i.e. a stochastic differential equation that captures the trajectories of the weights $\mathbf{w}_i^k$. Using this connection, we prove existence and uniqueness of weak solutions of Eq. (3.9). In the proofs, we will often assume $\nu_0 = 1$, which amounts to a rescaling of time t .

For $\tau = 0$, the evolution defined by Eq. (3.9) corresponds to the gradient flow in Wasserstein metric for the risk function $R^\delta(\rho)$. For $\tau > 0$, it is the gradient flow for the free energy functional $F^\delta(\rho)$ defined below

$$F^\delta(\rho) = \frac{1}{2} R^\delta(\rho) - \tau S(\rho), \quad S(\rho) = - \int \rho(\mathbf{w}) \log \rho(\mathbf{w}) d\mathbf{w}. \quad (3.11)$$

3.5 Limit PDE, $\delta = 0$

As mentioned above, in the limit $\delta \rightarrow 0$ the risk function $R^\delta(\rho)$ is well approximated by $R : \mathcal{L}^2(\Omega) \rightarrow \mathbb{R}$, where $R(\rho) = \nu_0 \|f - \rho\|_{\mathcal{L}^2(\Omega)}^2$, cf. Eq. (1.7).

The corresponding Wasserstein gradient flow is also known as *viscous porous medium equation* [Váz07] and it is given by

$$\partial_t \rho_t(\mathbf{w}) = -\nu_0 \nabla \cdot (\rho_t(\mathbf{w}) \nabla f(\mathbf{w})) + \frac{\nu_0}{2} \Delta(\rho_t^2(\mathbf{w})) + \tau \Delta \rho_t(\mathbf{w}), \quad (3.12)$$

with initial and boundary conditions

$$\begin{aligned} \rho_0 &= \rho_{\text{init}}, \\ \langle \mathbf{n}(\mathbf{w}), \nu_0 \rho_t(\mathbf{w}) \nabla(f(\mathbf{w}) - \rho_t(\mathbf{w})) - \tau \nabla \rho_t(\mathbf{w}) \rangle &= 0 \quad \forall \mathbf{w} \in \partial\Omega. \end{aligned} \quad (3.13)$$

In Appendix A, we give the definition of a weak solution for the PDE (3.12) with initial and boundary conditions (3.13). We also prove that the weak solution of the PDE (3.12) is unique, under a mild integrability condition. Again, in proofs we will assume without loss of generality $\nu_0 = 1$.

As in the $\delta > 0$ case, the evolution defined by Eq. (3.12) is the gradient flow for the free energy $F(\rho) = (1/2)R(\rho) - \tau S(\rho)$. Our analysis uses a key property of the risk function $R(\rho) = \nu_0 \|f - \rho\|_{\mathcal{L}^2(\Omega)}^2$ (and the free energy): displacement convexity [McC97]. For the reader's convenience, we recall its definition here, referring to [AGS08, Vil08, San15] for further background. Given two probability measures $\rho_0, \rho_1 \in \mathcal{P}_2(\Omega)$, their W_2 distance is defined by

$$W_2(\rho_0, \rho_1)^2 = \inf_{\gamma \in \Gamma(\rho_0, \rho_1)} \int \|\mathbf{x} - \mathbf{y}\|_2^2 \gamma(d\mathbf{x}, d\mathbf{y}), \quad (3.14)$$

where the infimum is taken over the set $\Gamma(\rho_0, \rho_1)$ of couplings of ρ_0, ρ_1 (i.e. probability measures on $\Omega \times \Omega$ whose first marginal coincides with ρ_0 , and second with ρ_1). The infimum is achieved by weak compactness of $\mathcal{P}_2(\Omega)$.

The metric space $(\mathcal{P}_2(\Omega), W_2)$ is a ‘length space,’ and in particular it is possible to construct geodesics, i.e. paths of minimum length connecting any two probability measures ρ_0, ρ_1 . Geodesics have a simple description. Let γ_* be the coupling achieving the infimum in the definition of $W_2(\rho_0, \rho_1)$. Letting $(\mathbf{X}_0, \mathbf{X}_1) \sim \gamma_*$, we define ρ_t to be the distribution of $\mathbf{X}_t = (1-t)\mathbf{X}_0 + t\mathbf{X}_1$. The curve $t \mapsto \rho_t$, indexed by $t \in [0, 1]$ turns out to be the geodesic between ρ_0 and ρ_1 in $(\mathcal{P}_2(\Omega), W_2)$.

Displacement convexity is convexity along geodesics. Namely, a function $\mathcal{F} : \mathcal{P}_2(\Omega) \rightarrow \mathbb{R}$ is λ -strongly displacement convex if

$$(1-t)\mathcal{F}(\rho_0) + t\mathcal{F}(\rho_1) - \mathcal{F}(\rho_t) \geq \frac{1}{2}\lambda t(1-t)W_2(\rho_0, \rho_1)^2. \quad (3.15)$$

A useful observation is that displacement convexity implies that all local minima of $\mathcal{F}$ are global minimizers. Indeed, by (3.15) it is straightforward to see that $\mathcal{F}$ has at most one global minimizer ρ_* . Also, for every other point ρ , the geodesic between ρ and ρ_* is a strictly decreasing path for the function $\mathcal{F}$. Now, suppose that $\bar{\rho} \neq \rho_*$ is a local minimum. Then, there exists a neighborhood U around $\bar{\rho}$ such that, for any $\rho \in U$, $\mathcal{F}(\rho) \geq \mathcal{F}(\bar{\rho})$. However, the strictly decreasing path between $\bar{\rho}$ and ρ_* passes through the neighborhood U , which leads to a contradiction and so $\rho = \rho_*$.

It follows from [McC97] that the risk function $R(\rho)$ and the free energy $F(\rho)$ are strongly displacement convex.

Remark 3.2. The concavity assumption on the regression function f (Assumption (A2)) defines a nonparametric class under which global convergence can be established, with convergence rates

uniquely determined by the curvature α (in the limit $N \rightarrow \infty$, $\delta \rightarrow 0$). Nonparametric estimation of concave functions has attracted considerable attention over recent years, see e.g. [HD13, CS16], and is –by itself– an interesting domain of applicability.

However, our projected SGD algorithm is potentially applicable to any data set, and will return a meaningful estimate $\hat{f}$ regardless whether f is concave or not. Indeed, in the next section we present numerical simulations indicating convergence to a near-global optimum even for non-concave functions f .

From mathematical point of view, Assumption (A2) is only used to show the convergence of the solution of the viscous porous medium equation (limit PDE, $\delta = 0$) to the unique global minimizer of the free energy $F(\rho) = (1/2)R(\rho) - \tau S(\rho)$, as formally stated in Theorem F.8. Concavity is not needed for the other results in the paper, namely approximating the SGD trajectory with the solution of the PDE ($\delta > 0$), see Theorem 5.1, and the convergence of the solution of the PDE ($\delta > 0$) to the solution of the viscous porous medium equation, see Theorem 5.2. It is therefore foreseeable a more general analysis that relaxes the concavity assumption.

4 Numerical illustrations

In this section we provide some simple numerical illustrations of our setting, and compare numerical results with the predictions of the Wasserstein gradient flow theory.

It is easy to construct examples of strongly concave functions, satisfying our assumptions. One can start from any strongly concave continuous function f_0 on a compact convex set Ω , add a constant to make it non-negative, and multiply it by a constant to normalize its integral. The resulting function $f(\mathbf{x}) = (c_1 + f_0(\mathbf{x}))/c_2$ satisfies our conditions. Notable examples of concave functions are given by log-moment generating functions $f_0(\mathbf{x}) = -\log \mathbb{E}_{\mathbf{Z}} \exp\{\langle \mathbf{x}, \mathbf{Z} \rangle\}$, where the random variable $\mathbf{Z}$ satisfies mild assumptions (e.g., it is bounded and its distribution is not supported on a proper subspace of $\mathbb{R}^d$). In general, given any twice differentiable function g_0 , the function $f_0(\mathbf{x}) = g_0(\mathbf{x}) - c_* \|\mathbf{x}\|_2^2$ is strongly concave for c_* large enough.

4.1 A one-dimensional concave function

We set $\Omega = [-1, 1]$ and $f(x) = (1 - e^{x-1})/(1 - e^{-2})$ (we choose the normalization so that $\int_{-1}^1 f(x) dx = 1$). Note that f is uniformly concave in $[-1, 1]$. We set the kernel K as follows:

$$K(x) = C_d \kappa(|x|), \quad \kappa(t) = \begin{cases} 1 - t^2 - 2t^3 + 2t^4 & \text{for } t \leq c_0 = 1, \\ 0 & \text{otherwise,} \end{cases} \quad (4.1)$$

where C_d is a normalization constant ensuring that $\int_{-1}^1 K(x) dx = 1$. The initialization ρ_{init} is a truncated Gaussian: $\rho_{\text{init}}(x) = c \cdot \exp(-x^2/(2\sigma^2)) \mathbf{1}_{[-1,1]}(x)$, with $\sigma = 1/3$.

We find empirically that standard stochastic gradient descent (SGD) without the projection $\mathbf{P}$ onto Ω^δ works well in this example, and consider this algorithm for simplicity in our first illustrations. We pick $N = 200$, $\tau = 0$ (noiseless SGD), and constant step size $\varepsilon = 10^{-6}$. In Figure 1, left column, we plot the true function $f(\cdot)$ together with the neural network estimate $\hat{f}(\cdot; \mathbf{w}^k)$ at several points in time t (time is related to the number of iterations k via $t = k\varepsilon$). Different

plots correspond to different values of δ with $\delta \in \{1/5, 1/10, 1/20\}$. We observe that the network estimates $\hat{f}(\cdot; \mathbf{w}^k)$ seem to converge to a limit curve which is an approximation of the true function f . As expected, the quality of the approximation improves as δ gets smaller.

In the right column, we report the evolution of the population risk (1.2) normalized by $\|f\|_{\mathcal{L}^2(\Omega)}^2$. For comparison, we plot the evolution of the risk (1.7) as predicted by the limit PDE (3.12) with $\tau = 0$. We solve the PDE (3.12) numerically using a finite difference scheme that enforces the conservation law $\int \rho(x, t) dx = 1$, see, e.g., [Tho13]. In the finite difference scheme, we choose time step and spatial step $\Delta t = 10^{-5}$ and $\Delta x = 10^{-2}$, respectively. The curve obtained by this numerical solution appears to capture well the evolution of SGD towards optimality. The main difference is that, while the PDE (3.12) corresponds to $\delta = 0$, and hence evolves towards a global optimum at zero risk, SGD converges to a non-zero risk value, which can be interpreted as the approximation error, decreasing with δ .

In Figure 2, we illustrate the numerical solution of the PDE (3.12) by plotting (i) the regression function f together with the PDE solution ρ_t (which coincides with the prediction $\hat{f}$ at $\delta = 0$) at several times t , and (ii) the PDE prediction for the risk $R(\rho_t)$ (1.7) normalized with respect to $\|f\|_{\mathcal{L}^2(\Omega)}^2$ (this plot aggregates data from Figs. 1.(b), (d), (f)). We also compare the risk (1.7) to the population risk $R_N(\mathbf{w}^k)$ achieved by SGD for different values of δ . Note that, as δ becomes smaller, the risk converges to the predicted curve. The risk of the limit PDE (3.12) converges to 0 exponentially fast in t , as predicted by the strong displacement convexity of $R(\rho)$.

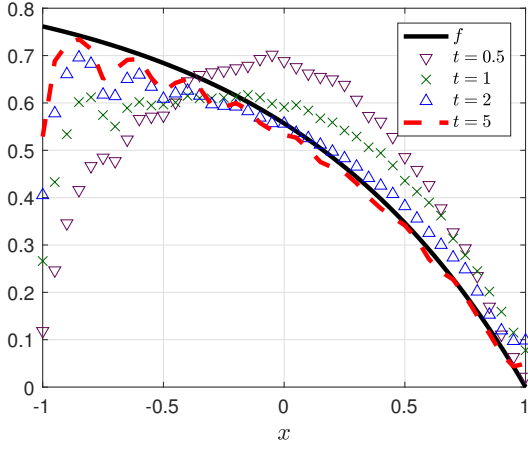
In Figure 3, we consider the SGD algorithm with projection $\mathbf{P}$, see (3.5). We pick $N = 200$, $\tau = 0$, $\varepsilon = 10^{-6}$ and $\delta = 1/20$. On the left, we illustrate the evolution of the value of 40 weights chosen at random; and on the right, we plot the histogram of their empirical distribution at $t = 5$. Note that this histogram matches well the regression function f plotted in black.

4.2 A two-dimensional concave example

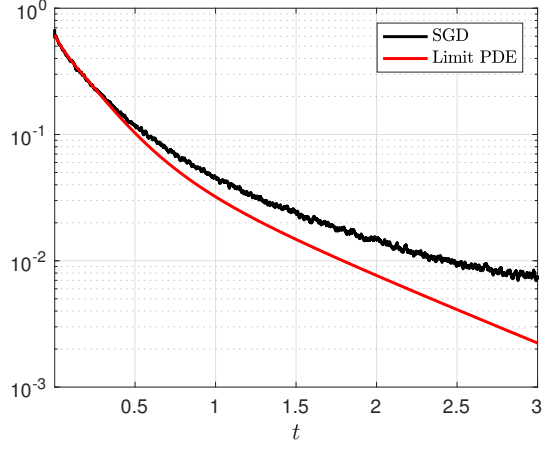
Next, we consider a two-dimensional example. We set $\Omega = [-1, 1]^2$ and

$$f(\mathbf{x}) = \frac{c_1 - \log(e^{\langle \mathbf{q}_1, \mathbf{x} \rangle} + e^{\langle \mathbf{q}_2, \mathbf{x} \rangle})}{c_2}, \quad (4.2)$$

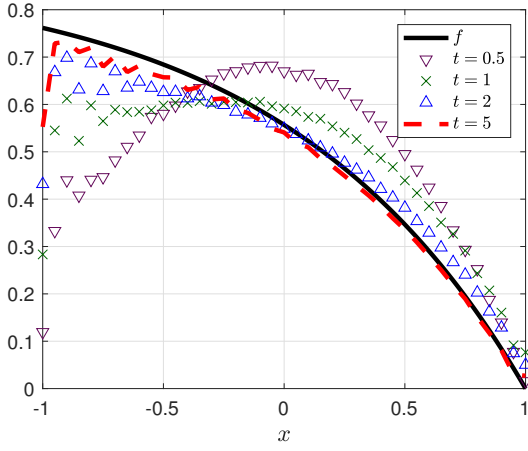
with $\mathbf{q}_1 = (2.5127, -2.4490)$, $\mathbf{q}_2 = (0.0596, 1.9908)$ and where c_1 and c_2 are chosen so that f is non-negative and $\int_{\Omega} f(\mathbf{x}) d\mathbf{x} = 1$. The kernel K is given by $K(\mathbf{x}) = C_d \kappa(|\mathbf{x}|)$, where κ is defined in (4.1) and C_d is a normalization constant ensuring that $\int_{\mathbf{B}(\mathbf{0}; 1)} K(\mathbf{x}) d\mathbf{x} = 1$. Again, the initialization ρ_{init} is a truncated Gaussian: $\rho_{\text{init}}(\mathbf{x}) = c \cdot \exp(-|\mathbf{x}|^2/(2\sigma^2)) \mathbf{1}_{[-1, 1]^2}(\mathbf{x})$, with $\sigma = 1/3$. We compare the normalized risk of SGD with no projection $\mathbf{P}$ ($N = 2000$, $\tau = 0$ and $\varepsilon = 10^{-6}$) for $\delta \in \{1/3, 1/5, 1/10\}$ with that of the limit PDE (3.12). Figure 4 shows that, already at $\delta = 1/10$, the risk of SGD converges to the predicted curve and the risk of the limit PDE (3.12) tends to 0 exponentially fast in t .



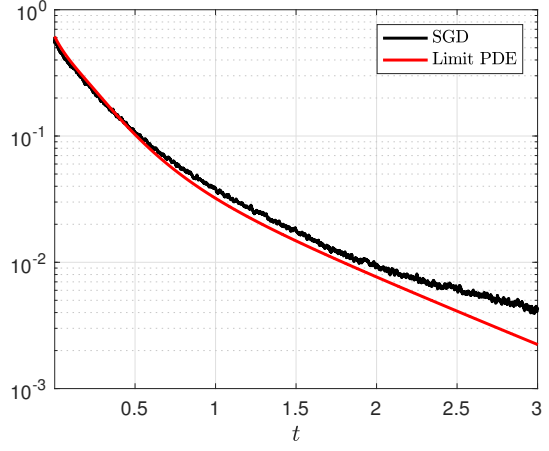
(a) Function f and SGD estimates, $\delta = 1/5$.



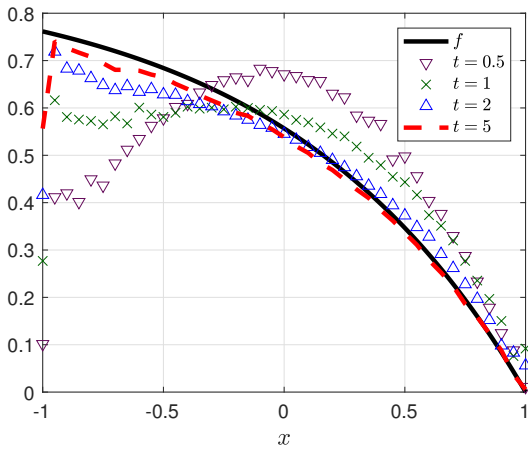
(b) Normalized risk, $\delta = 1/5$.



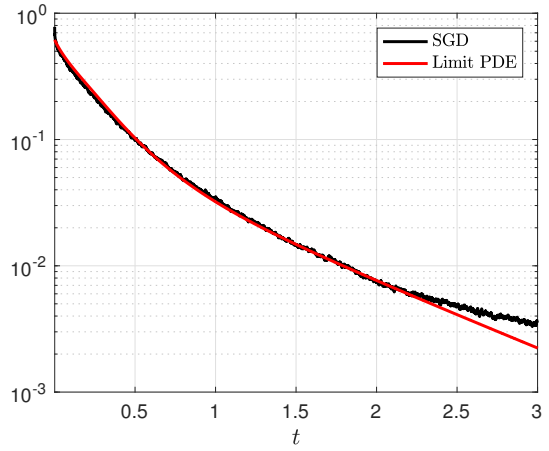
(c) Function f and SGD estimates, $\delta = 1/10$.



(d) Normalized risk, $\delta = 1/10$.



(e) Function f and SGD estimates, $\delta = 1/20$.



(f) Normalized risk, $\delta = 1/20$.

Figure 1: Dynamics of SGD update (3.5) at different times t and for different values of δ .

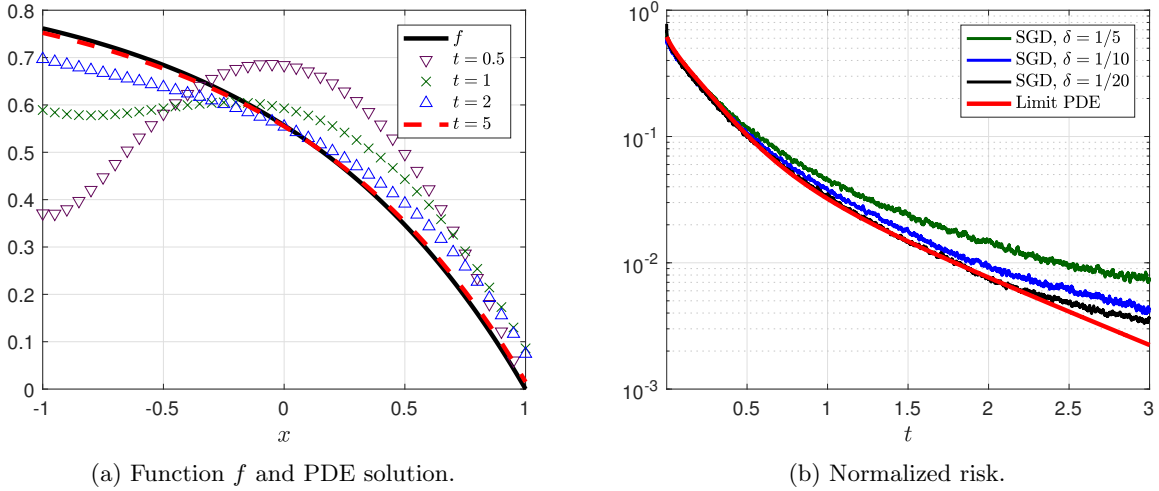


Figure 2: Dynamics of limit PDE (3.12) at different times t .

4.3 Comparing feature learning to random features

As discussed in the introduction, it is useful to consider the more general model

$$\hat{f}(\mathbf{x}; \mathbf{w}, \mathbf{a}) = \sum_{i=1}^N a_i K^\delta(\mathbf{x} - \mathbf{w}_i), \quad (4.3)$$

with parameters $\mathbf{a} = (a_1, \dots, a_N)$ as well as $\mathbf{w} = (\mathbf{w}_1, \dots, \mathbf{w}_N)$. This setting allows to compare two different approaches:

- (i) *Random feature regression*: the weights $\mathbf{w}$ are chosen independently of the labels y_i (we allow for dependence on the covariates $\mathbf{x}_i$).
- (ii) *Feature learning*: the weights $\mathbf{w}$ depend on the data $(y_i, \mathbf{x}_i)$.

In order to compare these two approaches, we assume to be given i.i.d. data $\{(y_i, \mathbf{x}_i)\}_{i \leq n}$, with $\mathbf{x}_i \sim \text{Unif}(\Omega)$, $y_i = f(\mathbf{x}_i)$ and determine the parameters $\mathbf{a}$ by the same method, ridge regression. More explicitly, define the matrix $\mathbf{Z} \in \mathbb{R}^{n \times N}$ as $(\mathbf{Z})_{i,j} = K^\delta(\mathbf{x}_i - \mathbf{w}_j)$. Then, we estimate $\mathbf{a}$ via

$$\hat{\mathbf{a}} = (\mathbf{Z}^\top \mathbf{Z} + \lambda \mathbf{I})^{-1} \mathbf{Z}^\top \mathbf{y}, \quad (4.4)$$

where λ is chosen via cross-validation on a hold-out set, comprising 10% of the samples.

In Figure 5, we compare the performance of three different ways to construct the weights $\mathbf{w}$: ‘*random $\mathbf{w}$* ,’ we choose the weights $\mathbf{w}_i$ independently and uniformly at random in Ω (blue triangles pointing down); ‘ *$\mathbf{w}$ = data points*,’ we choose the weights $\mathbf{w}_i$ uniformly at random among the data points (green circles); ‘*optimized $\mathbf{w}$* ,’ we use the output of the projected SGD algorithm of the previous sections (red triangles pointing up). The first two can be regarded as ‘random features’ approaches, while the latter is a ‘feature learning’ method.

For the optimized $\mathbf{w}$, we use exactly the same algorithm as in (3.5) (without coefficients $\mathbf{a}$ in the SGD update), with the only difference that each SGD step is carried out with respect

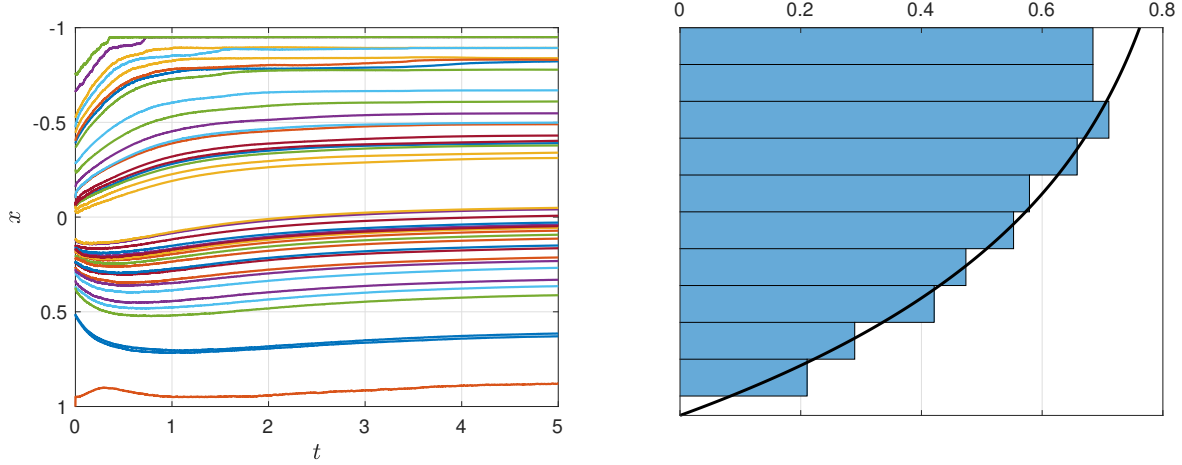


Figure 3: Evolution of the value of 40 weights chosen at random and histogram of their empirical distribution at time $t = 5$.

to an independent sample from the empirical data, with replacement. SGD is stopped after $k_{\max}$ iteration, and the coefficient $\hat{\mathbf{a}}$ are computed according to (4.3). Notice that this procedure is probably suboptimal, and it would be better to optimize $\mathbf{a}$ and $\mathbf{w}$ jointly: we choose this simpler two-stage procedure to have a more direct application of the algorithm analyzed in the paper, and a comparison with the random feature methods. We set $\tau = 0$ (noiseless SGD), and constant step size $\varepsilon = 5 \cdot 10^{-4}$. The number of iterations $k_{\max} \in \{5 \cdot 10^3, 15 \cdot 10^3, 5 \cdot 10^4, 15 \cdot 10^4, 5 \cdot 10^5, 15 \cdot 10^5\}$ is chosen via cross-validation, by using the same hold-out set employed to optimize λ .

We set $\Omega = [-1, 1]^4$ and define $y_j = f(\mathbf{x}_j)$, where $f(\mathbf{x})$ takes the form (4.2) with $\mathbf{q}_1 = (-0.3832, 0.3074, -0.3198, 0.4792)$ and $\mathbf{q}_2 = (0.3502, -0.1471, 0.1685, 0.0546)$. Again, c_1 and c_2 are chosen so that f is non-negative and $\int_{\Omega} f(\mathbf{x}) d\mathbf{x} = 1$; the kernel K is given by $K(\mathbf{x}) = C_d \kappa(|\mathbf{x}|)$, where κ is defined in Eq. (4.1) and C_d ensures that $\int_{\mathbf{B}(0;1)} K(\mathbf{x}) d\mathbf{x} = 1$.

After estimating $\mathbf{w}_i$ and a_i by either methods, we generate a test set of 10,000 samples and use it to estimate the generalization error. We perform 20 independent trials of the experiment, and we plot the average risk normalized by $\|f\|_{\mathcal{L}^2(\Omega)}^2$ together with the error bar at 1 standard deviation. In Figure 5-(a), we fix the number of neurons $N = 200$ and we plot the normalized risk as a function of the number of data points n . In Figure 5-(b), we fix the number of samples n to 2000 and we plot the normalized risk as a function of the number of neurons N . The data set used for cross-validation has size $\max(n/10, 40)$. Note that feature learning leads to improved performance in both settings. The improvement becomes more pronounced with the sample size n , presumably because a better set of weights $\mathbf{w}_i$ can be learnt. On the other hand, when the number of neurons N becomes very large, random $\mathbf{w}_i$'s are already covering Ω densely enough, and there is no significant advantage in feature learning.

4.4 A non-concave one-dimensional example

We set $\Omega = [-1, 1]$ and $f(x) = (x + \sin(5x - \pi/2) - c_1)/c_2$, where c_1 and c_2 are chosen so that f is non-negative and $\int_{\Omega} f(x) dx = 1$. Note that the target function f is bimodal, thus it is not concave. We perform the same numerical experiment described in Section 4.1. In Figure 6, left column, we

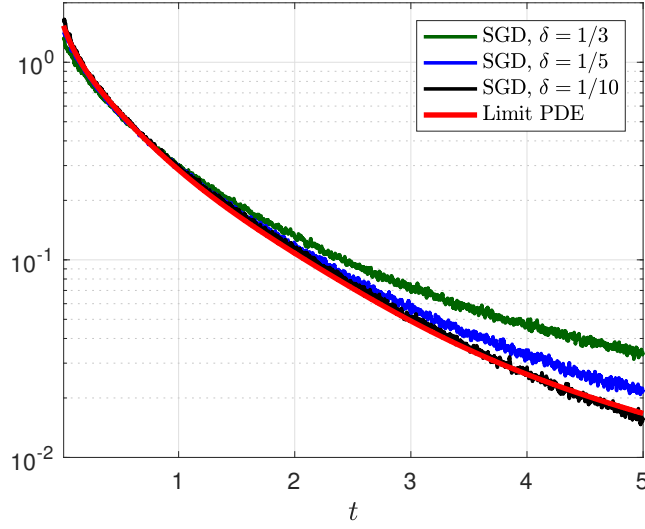
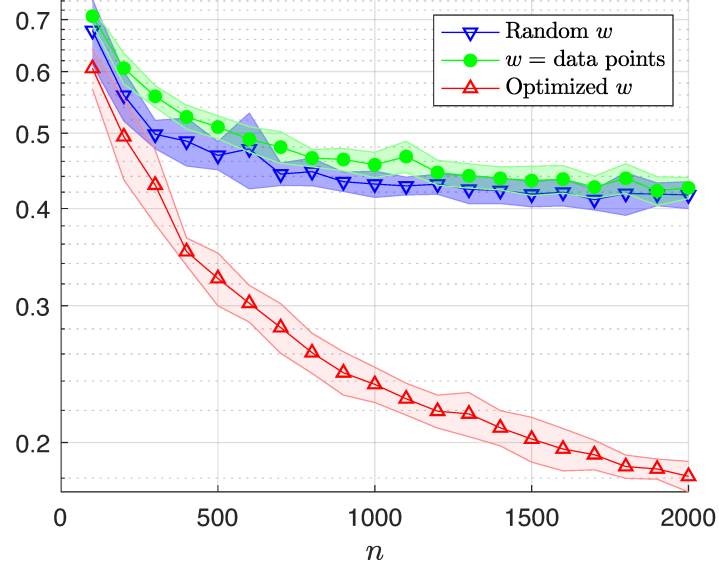


Figure 4: Normalized risk of SGD for different values of δ compared with that of the limit PDE for a two-dimensional example.

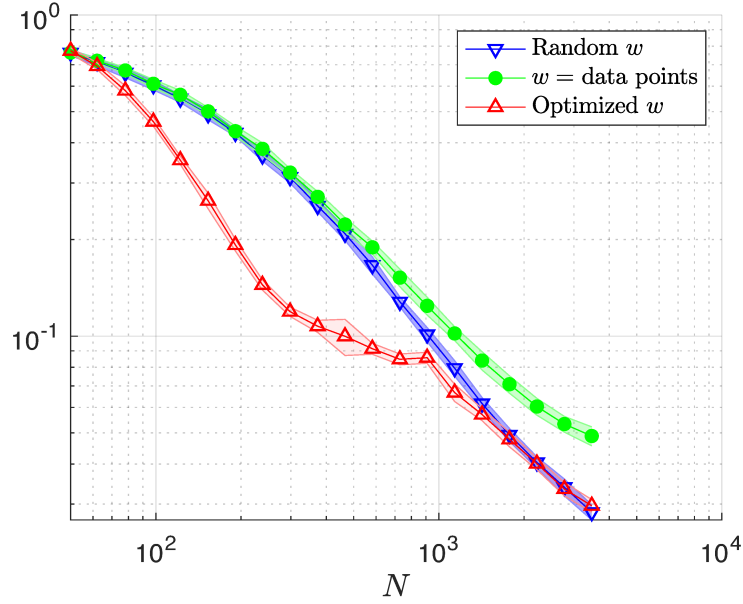
plot the true function $f(\cdot)$ together with the neural network estimate $\hat{f}(\cdot; \mathbf{w}^k)$ at several points in time t , where different plots correspond to different values of $\delta \in \{1/5, 1/10, 1/20\}$. In the right column, we report the evolution of the population risk (1.2) normalized by $\|f\|_{\mathcal{L}^2(\Omega)}^2$. In Figure 7, we plot (i) the regression function f together with the PDE solution ρ_t at several times t , and (ii) the PDE prediction for the risk $R(\rho_t)$ (1.7) (normalized with respect to $\|f\|_{\mathcal{L}^2(\Omega)}^2$) compared with the population risk $R_N(\mathbf{w}^k)$ achieved by SGD for different values of δ . Even if the target function is not concave, the results are similar to those presented in the concave case: (i) the network estimates $\hat{f}(\cdot; \mathbf{w}^k)$ seem to converge to a limit curve which is an approximation of the true function f , (ii) the quality of the approximation improves as δ gets smaller, and (iii) the risk of the limit PDE (3.12) converges to 0 exponentially fast in t .

4.5 Failure for small N

We repeat the same experiment described in Section 4.1 for a smaller number of neurons $N = 20$. As can be seen in Figures 8 and 9, the quality of the approximation becomes worse as δ gets smaller. This is expected because with small number of activations, reducing their bandwidth δ leads to a worse performance as they are all zero on a large part of the space. Put differently, the number of neurons is too small to guarantee convergence of SGD to the predictions of the Wasserstein gradient flow theory.

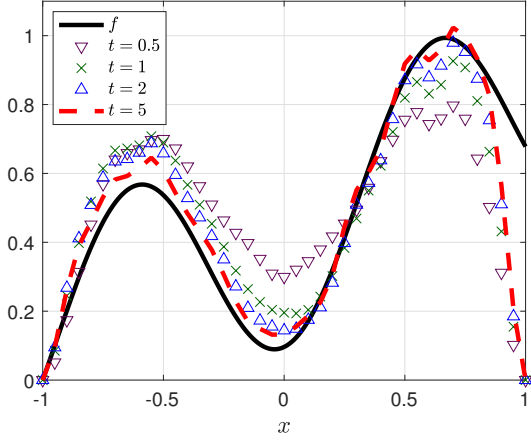


(a) $N = 200$, n varies on the x -axis.

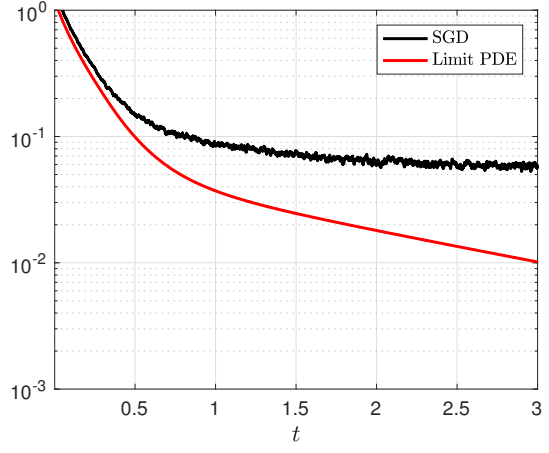


(b) $n = 2000$, N varies on the x -axis.

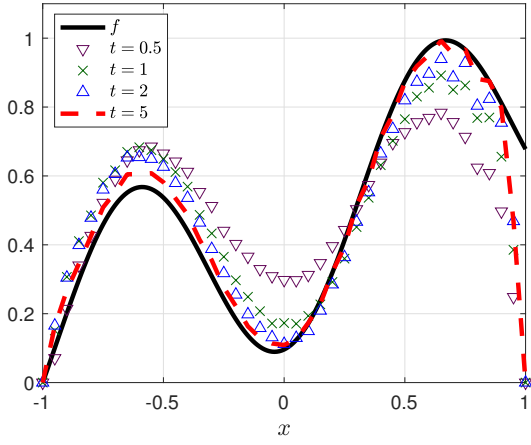
Figure 5: Generalization error achieved by fitting $\mathbf{a}$ from the data for three different choices of the weights $\mathbf{w}$: in red, the $\mathbf{w}_i$'s are optimized before-hand via SGD, as suggested in this paper; in blue, the $\mathbf{w}_i$'s are uniform in Ω ; and in green, the $\mathbf{w}_i$'s are equal to random data points.



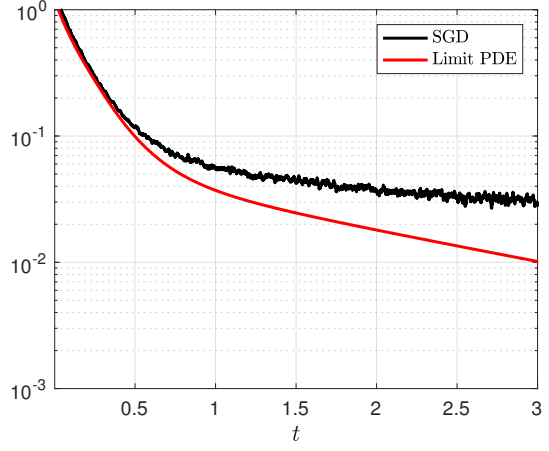
(a) Function f and SGD estimates, $\delta = 1/5$.



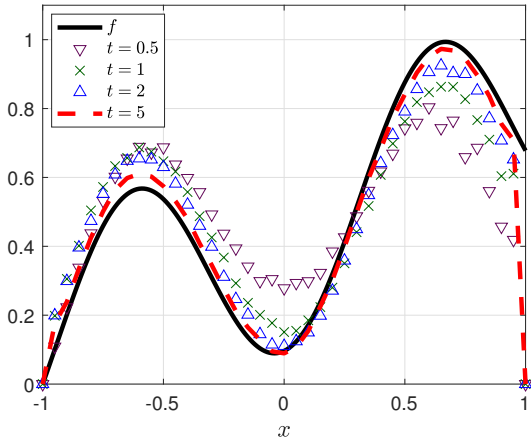
(b) Normalized risk, $\delta = 1/5$.



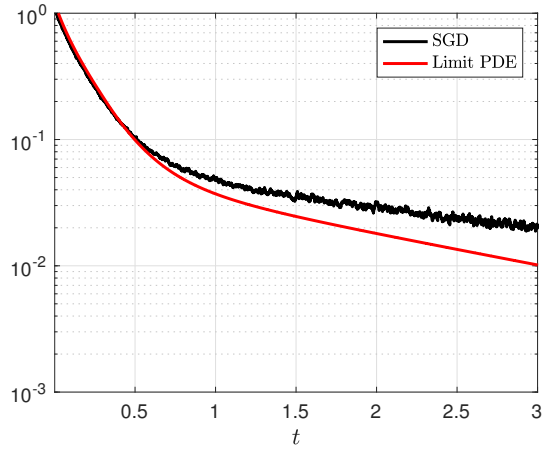
(c) Function f and SGD estimates, $\delta = 1/10$.



(d) Normalized risk, $\delta = 1/10$.

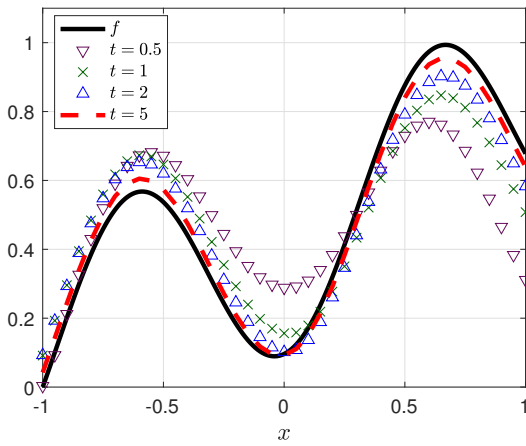


(e) Function f and SGD estimates, $\delta = 1/20$.

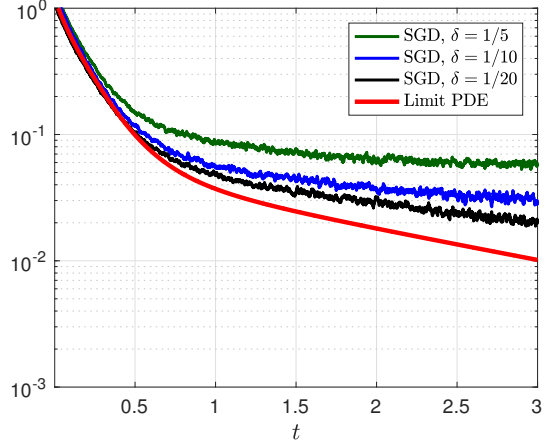


(f) Normalized risk, $\delta = 1/20$.

Figure 6: Dynamics of SGD update (3.5) at different times t and for different values of δ for a non-concave target function f .



(a) Function f and PDE solution.



(b) Normalized risk.

Figure 7: Dynamics of limit PDE (3.12) at different times t for a non-concave target function f .

5 Main results

5.1 Convergence of SGD to the PDE (3.9) at $\delta > 0$ fixed

We now state our result concerning the convergence of the SGD dynamics (3.5) to the PDE (3.9). Note that this result does not require concavity of f . Its proof is presented in Appendix D.

Theorem 5.1. *Assume that conditions (A1), (A3)-(A5) hold. Consider the SGD update (3.5) with initialization $(\mathbf{w}_i^0)_{i \leq N} \sim \text{i.i.d. } \rho_{\text{init}}^\delta$ and constant step size ε . For $t \geq 0$, let ρ_t be the unique solution of the PDE (3.9) with initial and boundary conditions (3.10), and assume $\text{supp}(\rho_{\text{init}}^\delta) \subseteq \mathbf{B}(\mathbf{0}, r)$. Then, for any fixed $t \geq 0$, $\rho_{\lfloor t/\varepsilon \rfloor}^{(N)} \Rightarrow \rho_t$ almost surely along any sequence $(N, \varepsilon = \varepsilon_N)$ such that $N \rightarrow \infty$, $\varepsilon_N \rightarrow 0$.*

Furthermore, for any $\delta \leq 1$, $T \geq 1$, $\varepsilon \leq 1$, $p \in \mathbb{N}$, and for any $g : \mathbb{R}^d \rightarrow \mathbb{R}$ with $\|g\|_{\text{Lip}} \leq 1$, the following happens with probability at least $1 - z^{-2p}$,

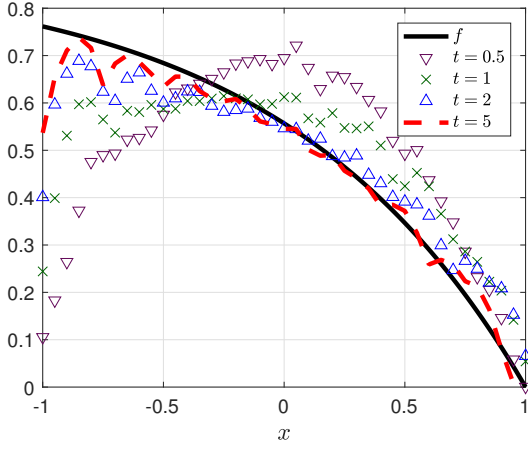
$$\begin{aligned} \sup_{k \in [0, T/\varepsilon] \cap \mathbb{N}} \left| \sum_{i=1}^N g(\mathbf{w}_i^k) - \int g(\mathbf{w}) \rho_{k\varepsilon}(\mathrm{d}\mathbf{w}) \right| &\leq z \text{err}(N, d, \varepsilon, \delta) e^{C_* p \delta^{-(d+2)} T}, \\ \sup_{k \in [0, T/\varepsilon] \cap \mathbb{N}} |R_N(\mathbf{w}^k) - R^\delta(\rho_{k\varepsilon})| &\leq z \text{err}(N, d, \varepsilon, \delta) e^{C_* p \delta^{-(d+2)} T}, \end{aligned} \quad (5.1)$$

where

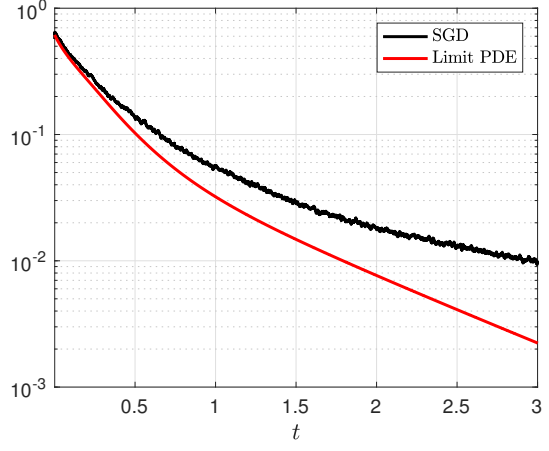
$$\text{err}(N, d, \varepsilon, \delta) = \sqrt{\frac{d}{N}} \vee (\delta^{-2d-1} r (d^2 \varepsilon \log(1/\varepsilon))^{1/4}). \quad (5.2)$$

Our proof is based on the same approach developed in [MMN18]. We prove that solutions of the PDE (3.9) are in correspondence with distributions over trajectories $(\mathbf{X}_t)_{t \geq 0}$ in Ω satisfying the following stochastic differential equation

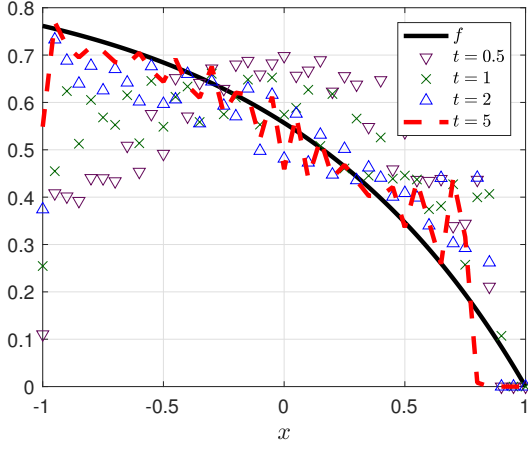
$$\mathrm{d}\mathbf{X}_t = -\nabla \Psi(\mathbf{X}_t, \rho_t) \mathrm{d}t + \sqrt{2\tau} \mathrm{d}\mathbf{B}_t + \mathrm{d}\Phi_t, \quad (5.3)$$



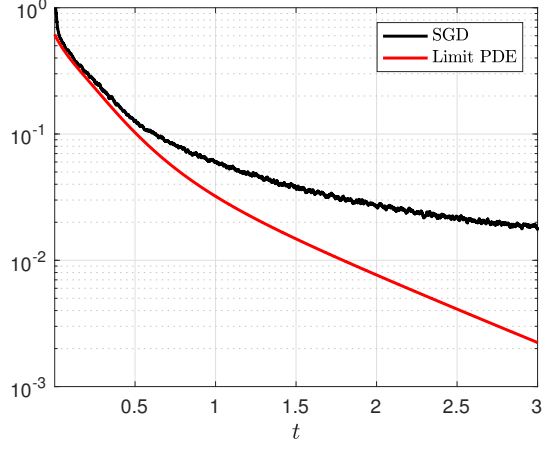
(a) Function f and SGD estimates, $\delta = 1/5$.



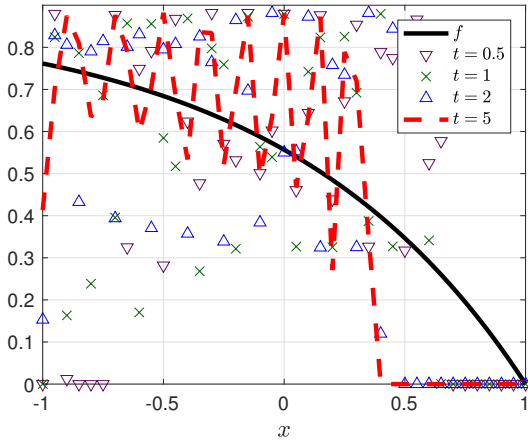
(b) Normalized risk, $\delta = 1/5$.



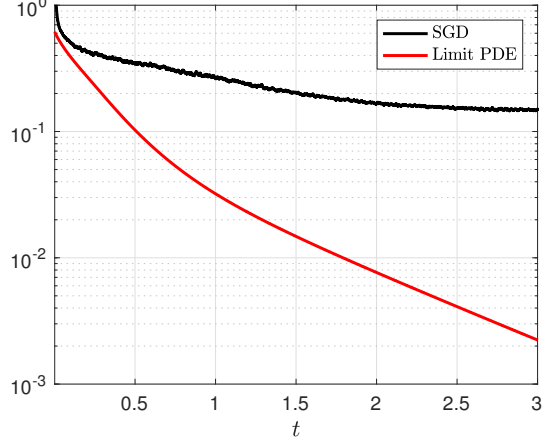
(c) Function f and SGD estimates, $\delta = 1/10$.



(d) Normalized risk, $\delta = 1/10$.



(e) Function f and SGD estimates, $\delta = 1/20$.



(f) Normalized risk, $\delta = 1/20$.

Figure 8: Dynamics of SGD update (3.5) at different times t and for different values of δ when the number of neurons is too small ($N = 20$).

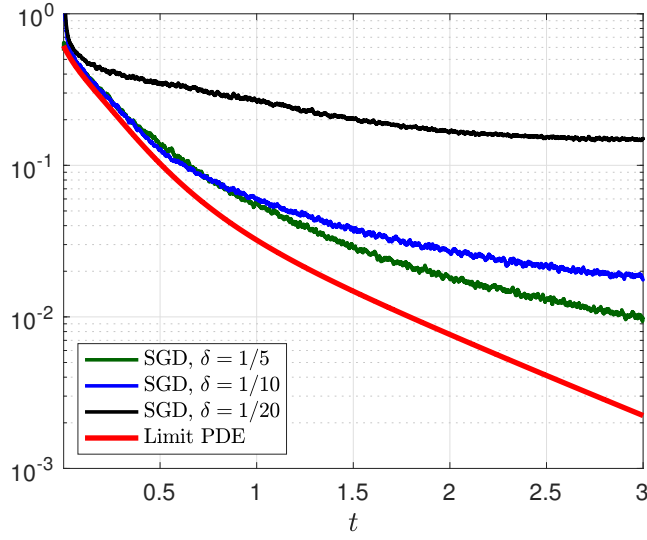


Figure 9: Normalized risk of the limit PDE (3.12) and of the SGD update (3.5) when the number of neurons is too small ($N = 20$).

where $(\mathbf{B}_t)_{t \geq 0}$ is a standard Brownian motion and $d\Phi_t$ is the boundary reflection (in the sense of a Skorokhod problem). The density ρ_t is determined, self consistently, via $\rho_t = \text{Law}(\mathbf{X}_t)$. We prove existence and uniqueness of solutions to this problem, and refer to the corresponding stochastic process $(\mathbf{X}_t)_{t \geq 0}$ as *nonlinear dynamics*. This in turn implies existence and uniqueness of the solutions of the PDE (3.9).

We next construct a coupling between the network weights $(\mathbf{w}_1^k, \dots, \mathbf{w}_N^k) \in (\Omega^\delta)^N$, and N i.i.d. trajectories of the nonlinear dynamics $(\mathbf{X}_1^t, \dots, \mathbf{X}_N^t) \in (\Omega^\delta)^N$. Controlling the expected distance in this coupling yields Theorem 5.1.

Remark 5.1. The error term in Eq. (5.1) is completely analogous to the error in a similar theorem proved in [MMN18]. The constant δ^{-d} appearing here is obtained by bounding the Lipschitz constant of $\nabla \Psi(\mathbf{w}; \rho)$. As already mentioned, the main technical difficulty with respect to [MMN18] is posed by the Neumann (reflecting) boundary conditions. Indeed, even if we are given a solution of the PDE (3.9), existence and uniqueness of solutions of the Skorokhod problem (5.3) is a highly non-trivial fact first established in [Tan79, LS84]. As a consequence, while the main proof idea is similar to the one in [MMN18], its implementation is significantly different.

Remark 5.2. As discussed in Appendix D, our proof applies to a more general version of the PDE (3.9) and correspondingly of the SGD dynamics (3.5), where Ψ takes the form $\Psi(\mathbf{w}, \rho) = V(\mathbf{w}) + \int U(\mathbf{w}, \mathbf{w}') \rho(d\mathbf{w}')$, for $V : \Omega \rightarrow \mathbb{R}$, $U : \Omega \times \Omega \rightarrow \mathbb{R}$ two smooth functions. The SGD update (3.5) is generalized as in [MMN18], and Theorem 5.1 holds with the terms containing δ (i.e., δ^{-2d-1} and $\delta^{-(d+2)}$) replaced by a constant that depends uniquely on $\|\nabla V\|_{\mathcal{L}^\infty(\Omega)}$, $\|\nabla U\|_{\mathcal{L}^\infty(\Omega \times \Omega)}$, $\|\nabla^2 V\|_{\mathcal{L}^\infty(\Omega)}$, $\|\nabla^2 U\|_{\mathcal{L}^\infty(\Omega \times \Omega)}$.

5.2 Convergence to the solutions of porous medium equation

We next prove that the solution of the PDE (3.9) converges, as $\delta \rightarrow 0$, to the unique solution of the porous medium equation (3.12). As for Theorem 5.1, this result does not rely on the concavity assumption for f .

Theorem 5.2. *Assume that conditions (A1) and (A3)-(A5) hold. Denote by ρ^δ the unique solution of the PDE (3.9) with initial condition $\rho_0^\delta = \rho_{\text{init}}$. Then*

- (a) *The porous medium equation (3.12) admits a weak solution $\rho : (t, \mathbf{x}) \mapsto \rho_t(\mathbf{x})$ with initial and boundary conditions (3.13). Further, this solution is unique under the additional condition $\rho \in \mathcal{L}^4([0, T] \times \Omega)$.*
- (b) *For almost all $t \in [0, T]$, we have $\rho_t^\delta \rightarrow \rho_t$ in $\mathcal{L}^2(\Omega)$ as $\delta \rightarrow 0$.*

While this statement is very natural at a heuristic level, its proof is actually the bulk of our technical work. Similar approximation results have been proved in the past by Oelschläger, Philipowski, Figalli [Oel02, Phi07, FP08], but they do not apply directly to the present case unless $f = 0$ (also, we have to deal with different boundary conditions).

Our proof follows a classical compactness argument, generalizing the approach of [FP08]. Namely we consider the sequence of trajectories $(\rho_t^\delta)_{t \in [0, T]}$ indexed by the width δ . We prove that this family is bounded and equicontinuous in $\mathcal{C}([0, T], \mathcal{P}_2(\Omega))$, and hence admits converging subsequences $(\rho_t^{\delta_n})_{t \in [0, T]} \rightarrow (\rho_t)_{t \in [0, T]}$. We next prove that any such converging subsequence converges in $\mathcal{L}^2(\Omega \times [0, T])$ and that the limit is a weak solution of the porous medium equation (3.12). Unfortunately, uniqueness of weak solutions of the PME (3.12) is –to the best of our knowledge– an open problem. However, we generalize methods from [Oel02] to show that any subsequential limit is actually in $\mathcal{L}^4(\Omega \times [0, T])$, and prove that the weak solution is unique under this condition. This allows us to conclude that $(\rho_t^\delta)_{t \in [0, T]}$ converges to this unique weak solution $(\rho_t)_{t \in [0, T]}$.

5.3 Global convergence of SGD

Let us now state the main result of this paper: SGD converges to a model with nearly optimal risk.

Theorem 5.3. *Assume that conditions (A1)-(A5) hold, and recall that $\alpha > 0$ is the concavity parameter of the function f , i.e., $\langle \mathbf{y}, \nabla^2 f(\mathbf{x}) \mathbf{y} \rangle \leq -\alpha |\mathbf{y}|^2$ for all $\mathbf{x} \in \Omega$, $\mathbf{y} \in \mathbb{R}^d$.*

Consider the SGD update (3.5) with initialization $(\mathbf{w}_i^0)_{i \leq N} \sim_{\text{i.i.d.}} \rho_{\text{init}}$ and constant step size ε . Assume $\text{supp}(\rho_{\text{init}}) \subseteq \mathbf{B}(\mathbf{0}; r)$. Then, for any $k \leq T/\varepsilon$, the following holds with probability at least $1 - 1/z$,

$$R_N(\mathbf{w}^k) \leq R_N(\mathbf{w}^0) e^{-2\alpha k \varepsilon} + 2\tau \Delta'(k, \varepsilon, d) + \Delta(N, \varepsilon, T, d, \delta, z), \quad (5.4)$$

where

$$\Delta'(k, \varepsilon, d) = \log |\Omega| - (1 - e^{-2\alpha k \varepsilon}) S(f) - S(\rho_{\text{init}}) e^{-2\alpha k \varepsilon}, \quad (5.5)$$

$$\lim_{\delta \rightarrow 0} \lim_{N \rightarrow \infty, \varepsilon \rightarrow 0} \Delta(N, \varepsilon, T, d, \delta, z) = 0. \quad (5.6)$$

Remark 5.3. The error term $2\tau \Delta'(k, \varepsilon, d)$ in Eq. (5.4) is always non-negative. In fact, $\Delta'(k, \varepsilon, d) \geq 0$ as $S(\rho) \leq \log |\Omega|$ for any $\rho \in \mathcal{P}_2(\Omega)$. Furthermore, by applying Jensen's inequality, we have that, for any $\rho \in \mathcal{P}_2(\Omega)$,

$$S(\rho) = - \int \rho(\mathbf{x}) \log \rho(\mathbf{x}) \, d\mathbf{x} \geq - \log \int \rho(\mathbf{x})^2 \, d\mathbf{x} = -2 \log \|\rho\|_{\mathcal{L}^2(\Omega)},$$

which gives the following upper bound

$$\Delta'(k, \varepsilon, d) \leq \log |\Omega| + 2 \left| \log \|f\|_{\mathcal{L}^2(\Omega)} \right| + 2 \left| \log \|\rho_{\text{init}}\|_{\mathcal{L}^2(\Omega)} \right|.$$

Recall that τ controls the variance of the noise, which is added at each step of the SGD algorithm for technical purposes. Thus, we can take τ sufficiently small so that the term $2\tau \Delta'(k, \varepsilon, d)$ is arbitrarily small.

Remark 5.4. The proof of Theorem 5.3 provides a somewhat more explicit expression for the error term $\Delta(N, \varepsilon, T, d, \delta, z)$ in Eq. (5.4). Namely, for an arbitrary but fixed $p \in \mathbb{N}$,

$$\Delta(N, \varepsilon, T, d, \delta, z) = \Delta_1(N, \varepsilon, T, d, z) + \Delta_2(\delta, T, d), \quad (5.7)$$

$$\Delta_1(N, \varepsilon, T, d, z) = \left(\sqrt{\frac{d}{N}} \vee \left(r \delta^{-2d-1} (d^2 \varepsilon \log(1/\varepsilon))^{1/4} \right) \right) \quad (5.8)$$

$$\cdot \exp \left\{ \sqrt{2C_* \delta^{-(d+2)} T \log(z)} \right\} \\ \lim_{\delta \rightarrow 0} \Delta_2(\delta, T, d) = 0. \quad (5.9)$$

The term Δ_1 bounds the error due to describing the SGD dynamics using the PDE (3.9). It vanishes when $N \rightarrow \infty$, $\varepsilon \rightarrow 0$, under the stated conditions. The term Δ_2 captures the error due to approximating the PDE (3.9) with the porous medium equation (3.12). Finally, the term $e^{-2\alpha k \varepsilon}$ describes the convergence to equilibrium of the solution of the porous medium equation.

The proof of Theorem 5.3 is presented in Appendix F and relies crucially on regularity results for the PDE (3.9) which are established in Appendix E.

More specifically, the proof is based on three steps, which we spell out once more:

- (i) We approximate the dynamics of SGD by the PDE (3.9) at $\delta > 0$ fixed. In doing so, we incur an error Δ_1 which is controlled using Theorem 5.1.
- (ii) We approximate the solution ρ_t^δ of the PDE (3.9) at $\delta > 0$ using the solution ρ_t of the porous medium equation (3.12), as stated in Theorem 5.2.
- (iii) We use results from [CJM⁺01, CMV03, CMV06] to prove that the latter solution converges exponentially fast to the global optimum, with rate $O(e^{-2\alpha t})$.

Given Theorems 5.1, 5.2, and the results of [CJM⁺01, CMV03, CMV06], this proof is relatively direct. We emphasize that, unlike Theorems 5.1, 5.2, the proof Theorem 5.3 relies in a crucial way on our structural assumptions, namely the concavity of f , and the structure of the bump-like activation $K_\delta(\mathbf{x} - \mathbf{w}_i)$.

Remark 5.5. If we settle for the less ambitious goal of proving global convergence without the explicit dimension-independent rate $e^{-2\alpha k\varepsilon}$, and there are no boundary conditions ($\Omega = \mathbb{R}^d$), we can achieve this goal using [MMN18, Theorem 5]. This result guarantees convergence in a number of SGD steps that potentially depends on τ (the noise injected in SGD) as well as the dimensions d , and the width δ , but does not require to assume strong concavity of f . On the other hand, numerical experiments are consistent with the conclusion that rates are independent of these parameters, cf. e.g. Fig. 1 where dependence on δ is explored.

6 Discussion

It is instructive to compare the general strategy followed in this paper (and in related work, e.g. [MMN18, MMM19]) and the results we obtain, to a more classical approach in theoretical statistics. For the sake of clarity, we will abstract away most of the details of the present problem, and focus on the most important differences.

Consider a general setting in which we want to minimize the population risk $R(\mathbf{w}) = \mathbb{E}_{y,\mathbf{x}} L(\mathbf{w}; y, \mathbf{x})$, where L is a non-convex loss function and $\mathbf{w} \in \mathbb{R}^D$ are parameters (in our problem $\mathbf{w} = (\mathbf{w}_1, \dots, \mathbf{w}_N)$ are the first-layer weights and $D = dN$). We are given n i.i.d. samples $\{(y_j, \mathbf{x}_j)\}_{j \leq n}$.

A standard theoretical analysis of this problem uses empirical risk minimization. Namely, we define the empirical risk $\hat{R}_n(\mathbf{w}) = \hat{\mathbb{E}}_{y,\mathbf{x}} L(\mathbf{w}; y, \mathbf{x})$ (with $\hat{\mathbb{E}}_n$ denoting the empirical average), and compute the minimizer $\hat{\mathbf{w}}_n \in \arg \min_{\mathbf{w}} \hat{R}_n(\mathbf{w})$, for instance by gradient descent. Theoretical analysis proceeds –conceptually– in two steps. First, one proves that the empirical risk minimizer is a near-minimizer of the population risk. Namely

$$R(\hat{\mathbf{w}}_n) \leq \min_{\mathbf{w}} R(\mathbf{w}) + \text{err}(D, n). \quad (6.1)$$

This is normally proved through a uniform convergence argument to establish a bound $\sup_{\mathbf{w}} |\hat{R}_n(\mathbf{w}) - R(\mathbf{w})| \leq \text{err}(D, n)/2$. Here $\text{err}(D, n)$ is an error term that (hopefully) vanishes as $n \rightarrow \infty$ for D fixed. Second, one proves that gradient descent (with respect to the cost function $\hat{R}_n$) converges to a minimizer $\hat{\mathbf{w}}_n$. This is achieved by showing that, with high probability, the landscape $\mathbf{w} \mapsto \hat{R}_n(\mathbf{w})$ satisfies some strong conditions that guarantee convergence of gradient descent (or other algorithms). For instance, one desirable (although not sufficient) property is that $\hat{R}_n$ does not have local minima other than the global minima, provided that the sample size is large enough. A substantial literature applies this general scheme (with significant refinements) to a variety of non-convex problems in high-dimensional statistics, including phase retrieval, clustering, matrix completion, error-in-variables models, and so on. We refer to [MBM⁺18] for examples and a more detailed survey.

Unfortunately this approach runs into substantial difficulties when treating complex models such as multi-layer neural networks. We can name at least two sources of difficulties. First of all, the number of parameters D in the model is often comparable with the sample size n , and therefore uniform convergence of the empirical risk to population risk does not hold. For instance, in the present model, we could use a number of parameters $Nd \gtrsim n$: indeed, such an example is considered in Figure 5-(a), where $Nd = 800$ and $n \in \{100, \dots, 2000\}$. Of course this problem can be addressed by constraining other measures of complexity than the number of parameters [Bar98], but the common practice is not to add such regularizers in the training.

The second source of difficulties is that studying the risk landscape, and ruling out local minima is extremely difficult, even if we limit ourselves to the $n = \infty$ limit, i.e. the population risk $R(\mathbf{w})$. In two-layers neural networks, part of this difficulty is due to the fact that the risk (1.2) is invariant under permutations of the N neurons, and hence it has (generically) at least $N!$ global minima related by permutations, and a large number of saddle points connecting them.

The approach pursued in this paper builds on two simple remarks, which are connected to the previous difficulties:

- (i) Uniform convergence of the empirical risk $\hat{R}_n(\mathbf{w})$ to the population risk $R(\mathbf{w})$ is not necessary, nor it is necessary to control the random deviations of the whole landscape of the empirical risk. What is instead important is to control the landscape of the empirical risk along the trajectory of gradient descent from a given initialization.

A convenient way to implement this idea is to consider SGD in a one-pass setting in which each sample is used only once. In the limit of small step size, this converges to gradient flow with respect to $R(\mathbf{w})$.

- (ii) Absence of local minima in the population landscape $R(\mathbf{w})$ is not necessary either. What is instead important is absence of local minima along the gradient flow trajectory for $R(\mathbf{w})$ or, more precisely, the fact that the gradient flow trajectory converges to a global minimum.

These remarks suggest the following proof strategy. Let $\mathbf{w}(t)$ denote the gradient flow trajectory from a given initialization $\mathbf{w}(0) = \mathbf{w}_0$ (namely $\dot{\mathbf{w}}(t) = -\nabla R(\mathbf{w}(t))$), and $\mathbf{w}^k$ be the (random) parameters produced after k SGD steps. We first prove that gradient flow converges to a global optimum, possibly with explicit convergence rate $\Delta(t)$:

$$R(\mathbf{w}(t)) \leq \min_{\mathbf{w}} R(\mathbf{w}) + \Delta(t), \quad (6.2)$$

where $\Delta(t) \rightarrow 0$ as $t \rightarrow \infty$. We then show that the SGD trajectory, after k steps, is well approximated by the gradient flow for $R(\mathbf{w})$ provided the step size ε is small. For instance we might prove that there exists a numerical constant c_0 such that, for any $k\varepsilon \leq T$, with high probability

$$|R(\mathbf{w}^k) - R(\mathbf{w}(k\varepsilon))| \leq \varepsilon^{c_0} \text{err}(T). \quad (6.3)$$

The reader might recognize that the last estimate is analogous to the one obtained in Theorem 5.1, while the estimate 6.2 is what we obtain from displacement convexity (after taking the limit $\delta \rightarrow 0$ using Theorem 5.2). Putting the two estimates together, and recalling that we can run a total of n SGD steps (in the one-pass setting), we get

$$R(\hat{\mathbf{w}}) \leq \min_{\mathbf{w}} R(\mathbf{w}) + \Delta(n\varepsilon) + \varepsilon^{c_0} \text{err}(n\varepsilon), \quad (6.4)$$

where we set $\hat{\mathbf{w}} = \mathbf{w}^k$. The error is reminiscent of a bias-variance tradeoff: the first term is a bias due to early stopping; the second is instead the stochastic approximation error. We can now optimize n as to minimize this error. For instance, if $\Delta(t) = e^{-c_1 t}$, and $\text{err}(T) = e^{c_2 T}$, we can choose $\varepsilon \propto (\log n/n)$, yielding $R(\hat{\mathbf{w}}) \leq \min_{\mathbf{w}} R(\mathbf{w}) + C(\log n)^{c_0}/n^{c'}$ where $c' = c_0 c_1/(c_1 + c_2)$.

In summary, within the present approach, the generalization error is bounded via a tradeoff between the convergence rate of gradient flow in the population risk, and the error of approximating

the gradient flow by SGD. A side benefit of this proof strategy is that it guarantees the existence of an efficient algorithm to compute the weights $\hat{\mathbf{w}}$.

As mentioned, the above discussion omits several challenges that are posed by the model treated in this paper. Most notably: (1) We are trying to optimize N weight vectors $\mathbf{w}_1, \dots, \mathbf{w}_N \in \mathbb{R}^d$, but the loss only depends on the empirical distribution of these vectors $\hat{\rho}^{(N)} = N^{-1} \sum_{i=1}^N \delta_{\mathbf{w}_i}$. It is therefore natural to define a gradient flow in the space of probability distributions, which is nothing but the PDE (3.9). This also help addressing the challenge posed by the fact that, as N increases, the dimension of the parameter space increases and convergence to the population behavior might fail. We are embedding all the values of N in the space $\mathcal{P}(\mathbb{R}^d)$. (2) We cannot prove a bound of the form (6.2) for the original PDE (3.9) and have to approximate this by the porous medium equation (3.12).

Because of these additional challenges, our bounds are not nearly as neat as in Eqs. (6.2), 6.3 and depend on the additional parameters d, δ : in particular, the approximation by the porous medium equation in Theorem 5.2 is non-quantitative. We therefore refrain from optimizing the tradeoff between convergence rate of gradient flow, and error in stochastic approximation, which would result in suboptimal statistical guarantees, and defer this objective to future work.

Acknowledgements

A. Javanmard was partially supported by an Outlier Research in Business (iORB) grant from the USC Marshall School of Business, a Google Faculty Research award and the NSF CAREER award DMS-1844481. M. Mondelli was supported by an Early Postdoc.Mobility fellowship from the Swiss National Science Foundation and by the Simons Institute for the Theory of Computing. A. Montanari was partially supported by grants NSF DMS-1613091, CCF-1714305, IIS-1741162 and ONR N00014-18-1-2729. This work was carried out in part while the authors were visiting the Simons Institute for the Theory of Computing.

A Uniqueness of weak solutions of limit PDE ($\delta = 0$)

In this appendix, we prove that the limit PDE obtained for $\delta \rightarrow 0$, namely the porous medium equation (3.12) has at most one solution in $\mathcal{L}^4(\Omega \times [0, T])$. Existence of such solutions will follow from the results of Appendix F, and in particular from Lemma F.4.

For the sake of clarity, we repeat the definitions of Section 3.5. Let $\Omega \subseteq \mathbb{R}^d$ be a compact convex set with $\mathcal{C}^2$ boundary. We denote by $\mathcal{P}_2(\Omega)$ the space of probability measures on Ω endowed with Wasserstein's W_2 distance. Since Ω is compact, the induced topology is equivalent to weak convergence. We consider the following PDE:

$$\partial_t \rho_t(\mathbf{w}) = -\nu_0 \nabla \cdot (\rho_t(\mathbf{w}) \nabla f(\mathbf{w})) + \frac{\nu_0}{2} \Delta(\rho_t^2(\mathbf{w})) + \tau \Delta \rho_t(\mathbf{w}), \quad (\text{A.1})$$

with initial and boundary conditions

$$\begin{aligned} \rho_0 &= \rho_{\text{init}}, \\ \langle \mathbf{n}(\mathbf{w}), \nu_0 \rho_t(\mathbf{w}) \nabla(f(\mathbf{w}) - \rho_t(\mathbf{w})) - \tau \nabla \rho_t(\mathbf{w}) \rangle &= 0 \quad \forall \mathbf{w} \in \partial\Omega. \end{aligned} \quad (\text{A.2})$$

Throughout this appendix, we adopt the notation $\Phi(\rho) = \tau \rho + \nu_0 \rho^2/2$. Let us formally define the concept of weak solutions for the PDE (A.1).

For the next statement, it is useful to recall that $\mathcal{C}^{2,1}(\Omega \times [0, T])$ denotes the class of functions $f : \Omega \times [0, T] \rightarrow \mathbb{R}$ with continuous partial derivatives $D_t f(\mathbf{x}, t)$, $D_{\mathbf{x}}^{\alpha} f(\mathbf{x}, t)$ for all $\|\alpha\|_2 \leq 2$.

Definition A.1 (Weak solution of limit PDE). *We say that $\rho \in \mathcal{C}([0, T], \mathcal{P}_2(\Omega))$ is a weak solution of the PDE (A.1), with initial and boundary conditions (A.2) if*

1. ρ_t has density $\rho(\cdot, t)$ with respect to Lebesgue measure, and $\rho \in \mathcal{L}^2(\Omega \times [0, T])$.
2. For any test function $h \in \mathcal{C}^{2,1}(\Omega \times [0, T])$, satisfying $\langle \mathbf{n}(\mathbf{x}), \nabla h(\mathbf{x}, t) \rangle = 0$ for all $\mathbf{x} \in \partial\Omega, t \in [0, T]$, we have

$$\begin{aligned} & \int_{\Omega} h(\mathbf{x}, T) \rho(\mathbf{x}, T) d\mathbf{x} - \int_{\Omega} h(\mathbf{x}, 0) \rho_{\text{init}}(\mathbf{x}) d\mathbf{x} \\ &= \int_0^T \int_{\Omega} \left[\partial_t h(\mathbf{x}, t) + \nu_0 \langle \nabla f(\mathbf{x}), \nabla h(\mathbf{x}, t) \rangle \right] \rho(\mathbf{x}, t) d\mathbf{x} dt \\ &+ \int_0^T \int_{\Omega} \Delta h(\mathbf{x}, t) \Phi(\rho)(\mathbf{x}, t) d\mathbf{x} dt. \end{aligned} \quad (\text{A.3})$$

We now prove a uniqueness result, under a mild integrability condition.

Lemma A.2 (Uniqueness of limit PDE). *Let $\rho, \tilde{\rho} \in \mathcal{L}^4(\Omega \times [0, T])$ be two weak solutions of the PDE (A.1) with initial and boundary conditions (A.2), in the sense of Definition A.1. Then, $\rho = \tilde{\rho}$, almost everywhere.*

Proof. Note that setting $\nu_0 = 1$ corresponds to scaling time by a factor ν_0 and to substituting τ with $\tau \nu_0$. Since the proof holds for any $\tau > 0$, without loss of generality we can set $\nu_0 = 1$.

The proof follows ideas from [Váz07, Theorem 6.5]. We write the identity (A.3) for ρ and $\tilde{\rho}$ and subtract them to get

$$\begin{aligned} & \int_{\Omega} h_T(\mathbf{x}) (\rho_T(\mathbf{x}) - \tilde{\rho}_T(\mathbf{x})) d\mathbf{x} \\ &= \int_0^T \int_{\Omega} \left[\partial_t h_t(\mathbf{x}) + \langle \nabla f(\mathbf{x}), \nabla h_t(\mathbf{x}) \rangle \right] (\rho_t(\mathbf{x}) - \tilde{\rho}_t(\mathbf{x})) d\mathbf{x} dt \\ &+ \int_0^T \int_{\Omega} \Delta h_t(\mathbf{x}) (\Phi(\rho_t(\mathbf{x})) - \Phi(\tilde{\rho}_t(\mathbf{x}))) d\mathbf{x} dt, \end{aligned} \quad (\text{A.4})$$

where we use the shorthand $\rho_t(\mathbf{x}) \equiv \rho(\mathbf{x}, t)$ and $h_t(\mathbf{x}) \equiv h(\mathbf{x}, t)$. Define $u_t = \rho_t - \tilde{\rho}_t$ and $\eta_t = \tau + (\rho_t + \tilde{\rho}_t)/2$. Then,

$$\begin{aligned} \int_{\Omega} h_T(\mathbf{x}) u_T(\mathbf{x}) d\mathbf{x} &= \int_0^T \int_{\Omega} \left[\partial_t h_t(\mathbf{x}) + \langle \nabla f(\mathbf{x}), \nabla h_t(\mathbf{x}) \rangle \right] u_t(\mathbf{x}) d\mathbf{x} dt \\ &+ \int_0^T \int_{\Omega} \Delta h_t(\mathbf{x}) \eta_t(\mathbf{x}) u_t(\mathbf{x}) d\mathbf{x} dt. \end{aligned} \quad (\text{A.5})$$

Note that $\eta_t(\mathbf{x}) \geq \tau$ and define the truncated function $\eta_t^M = \min(M, \eta_t)$. We next choose a smooth test function $\theta : \Omega \times [0, T] \rightarrow \mathbb{R}_{\geq 0}$, $(\mathbf{x}, t) \mapsto \theta_t(\mathbf{x})$ and consider the following backward problem:

$$\begin{cases} \partial_t h_t(\mathbf{x}) + \hat{\eta}_t(\mathbf{x}) \Delta h_t(\mathbf{x}) + \langle \nabla f(\mathbf{x}), \nabla h_t(\mathbf{x}) \rangle + \theta_t(\mathbf{x}) = 0, & \forall \mathbf{x} \in \Omega, t \in [0, T], \\ \langle \mathbf{n}(\mathbf{x}), \nabla h_t(\mathbf{x}) \rangle = 0, & \forall \mathbf{x} \in \partial\Omega, t \in [0, T], \\ h_T(\mathbf{x}) = 0, & \forall \mathbf{x} \in \Omega. \end{cases} \quad (\text{A.6})$$

Here, $\hat{\eta}_t$ is a smooth approximation of η_t^M , such that $\tau \leq \hat{\eta}_t(\mathbf{x}) \leq M$. (We will make precise below in what sense $\hat{\eta}_t$ has to approximate η_t^M . For the moment, it can be a general smooth function satisfying the bounds $\tau \leq \hat{\eta}_t(\mathbf{x}) \leq M$.) Note that (A.6) is a backward parabolic problem with smooth coefficients and with Neumann boundary conditions. Hence, by classical results on quasilinear parabolic PDEs [LSU88], it admits a solution $h_t \in \mathcal{C}^{2,1}(\Omega \times [0, T])$. Rewriting (A.5) for such a test function h_t , we get

$$\int_0^T \int_{\Omega} \Delta h_t(\mathbf{x}) (\eta_t(\mathbf{x}) - \hat{\eta}_t(\mathbf{x})) u_t(\mathbf{x}) d\mathbf{x} dt = \int_0^T \int_{\Omega} \theta_t(\mathbf{x}) u_t(\mathbf{x}) d\mathbf{x} dt.$$

This immediately implies that

$$\int_0^T \int_{\Omega} \theta_t(\mathbf{x}) u_t(\mathbf{x}) d\mathbf{x} dt \leq \int_0^T \int_{\Omega} |\Delta h_t(\mathbf{x})| |\eta_t(\mathbf{x}) - \hat{\eta}_t(\mathbf{x})| |u_t(\mathbf{x})| d\mathbf{x} dt. \quad (\text{A.7})$$

By applying Cauchy-Schwarz inequality, we have that

$$\begin{aligned} & \int_0^T \int_{\Omega} |\Delta h_t(\mathbf{x})| |\eta_t(\mathbf{x}) - \hat{\eta}_t(\mathbf{x})| |u_t(\mathbf{x})| d\mathbf{x} dt \\ & \leq \left(\int_{\Omega \times [0, T]} \hat{\eta}_t(\mathbf{x}) (\Delta h_t(\mathbf{x}))^2 d\mathbf{x} dt \right)^{1/2} \\ & \quad \left(\int_{\Omega \times [0, T]} \frac{|\eta_t(\mathbf{x}) - \hat{\eta}_t(\mathbf{x})|^2}{\hat{\eta}_t(\mathbf{x})} u_t^2(\mathbf{x}) d\mathbf{x} dt \right)^{1/2}. \end{aligned} \quad (\text{A.8})$$

To bound the first term on the right-hand side of (A.8), we consider a smooth positive bounded function $\mu(t)$, defined on $[0, T]$, whose properties will be discussed later. Define the shorthand $\tilde{\theta}_t(\mathbf{x}) \equiv \theta_t(\mathbf{x}) + \langle \nabla f(\mathbf{x}), \nabla h_t(\mathbf{x}) \rangle$. We multiply the parabolic PDE (A.6) by $\mu(t)\Delta h_t(\mathbf{x})$ and integrate to obtain

$$\begin{aligned} \int_{\Omega \times [0, T]} \mu(t) \partial_t h_t(\mathbf{x}) \Delta h_t(\mathbf{x}) d\mathbf{x} dt &+ \int_{\Omega \times [0, T]} \mu(t) \hat{\eta}_t(\mathbf{x}) (\Delta h_t(\mathbf{x}))^2 d\mathbf{x} dt \\ &+ \int_{\Omega \times [0, T]} \mu(t) \tilde{\theta}_t(\mathbf{x}) \Delta h_t(\mathbf{x}) d\mathbf{x} dt = 0. \end{aligned} \quad (\text{A.9})$$

We next write

$$\begin{aligned} &\int_{\Omega \times [0, T]} \mu(t) \partial_t h_t(\mathbf{x}) \Delta h_t(\mathbf{x}) d\mathbf{x} dt \\ &\stackrel{(a)}{=} - \int_{\Omega \times [0, T]} \mu(t) \langle \nabla h_t(\mathbf{x}), \nabla (\partial_t h_t(\mathbf{x})) \rangle d\mathbf{x} dt \\ &= - \int_{\Omega \times [0, T]} \mu(t) \frac{1}{2} \frac{d}{dt} |\nabla h_t(\mathbf{x})|^2 d\mathbf{x} dt \\ &\stackrel{(b)}{=} \frac{1}{2} \int_{\Omega} \mu(0) |\nabla h_0(\mathbf{x})|^2 d\mathbf{x} - \frac{1}{2} \int_{\Omega} \mu(T) |\nabla h_T(\mathbf{x})|^2 d\mathbf{x} \\ &\quad + \frac{1}{2} \int_{\Omega \times [0, T]} \mu'(t) |\nabla h_t(\mathbf{x})|^2 d\mathbf{x} dt \\ &\stackrel{(c)}{\geq} \frac{1}{2} \int_{\Omega \times [0, T]} \mu'(t) |\nabla h_t(\mathbf{x})|^2 d\mathbf{x} dt \end{aligned} \quad (\text{A.10})$$

Here (a) follows from integration by parts in the integral over Ω and using the fact that $\langle \mathbf{n}(\mathbf{x}), \nabla h_t(\mathbf{x}) \rangle = 0$ for $\mathbf{x} \in \partial\Omega$ and $t \in [0, T]$. Also, (b) follows from integration by parts in the integral over t . Finally (c) holds because $h_T(\mathbf{x}) = 0$ for $\mathbf{x} \in \Omega$ and $\mu(0) \geq 0$.

Getting back to (A.9) and using the properties of function $\mu(t)$, we have

$$\begin{aligned} &\frac{1}{2} \int_{\Omega \times [0, T]} \mu'(t) |\nabla h_t(\mathbf{x})|^2 d\mathbf{x} dt + \int_{\Omega \times [0, T]} \mu(t) \hat{\eta}_t(\mathbf{x}) (\Delta h_t(\mathbf{x}))^2 d\mathbf{x} dt \\ &\leq - \int_{\Omega \times [0, T]} \mu(t) \tilde{\theta}_t(\mathbf{x}) \Delta h_t(\mathbf{x}) d\mathbf{x} dt \\ &= - \int_{\Omega \times [0, T]} \mu(t) \theta_t(\mathbf{x}) \Delta h_t(\mathbf{x}) d\mathbf{x} dt \\ &\quad - \int_{\Omega \times [0, T]} \mu(t) \langle \nabla f(\mathbf{x}), \nabla h_t(\mathbf{x}) \rangle \Delta h_t(\mathbf{x}) d\mathbf{x} dt \\ &= \int_{\Omega \times [0, T]} \mu(t) \langle \nabla \theta_t(\mathbf{x}), \nabla h_t(\mathbf{x}) \rangle d\mathbf{x} dt \\ &\quad - \int_{\Omega \times [0, T]} \mu(t) \langle \nabla f(\mathbf{x}), \nabla h_t(\mathbf{x}) \rangle \Delta h_t(\mathbf{x}) d\mathbf{x} dt \\ &\leq \int_{\Omega \times [0, T]} \mu(t) \langle \nabla \theta_t(\mathbf{x}), \nabla h_t(\mathbf{x}) \rangle d\mathbf{x} dt + \int_{\Omega \times [0, T]} \frac{\mu(t)}{2\tau} |\nabla f(\mathbf{x})|^2 |\nabla h_t(\mathbf{x})|^2 d\mathbf{x} dt \\ &\quad + \int_{\Omega \times [0, T]} \mu(t) \frac{\tau}{2} (\Delta h_t(\mathbf{x}))^2 d\mathbf{x} dt, \end{aligned} \quad (\text{A.11})$$

The penultimate step follows from integration by parts and the constraint $\langle \nabla h_t(\mathbf{x}), \mathbf{n}(\mathbf{x}) \rangle = 0$, for $\mathbf{x} \in \partial\Omega$ and $t \in [0, T]$, and the last step follows by applying Cauchy-Schwartz inequality. We continue by applying Cauchy-Schwartz inequality again to get

$$\begin{aligned} & \int_{\Omega \times [0, T]} \mu(t) \langle \nabla \theta_t(\mathbf{x}), \nabla h_t(\mathbf{x}) \rangle d\mathbf{x} dt \\ & \leq \int_{\Omega \times [0, T]} \left[\frac{\tau \mu(t)}{2C^2} |\nabla \theta_t(\mathbf{x})|^2 + \frac{C^2 \mu(t)}{2\tau} |\nabla h_t(\mathbf{x})|^2 \right] d\mathbf{x} dt, \end{aligned} \quad (\text{A.12})$$

where $C = \sup_{\mathbf{x} \in \Omega} |\nabla f(\mathbf{x})|$. Combining Equations (A.11) and (A.12), we get

$$\begin{aligned} & \frac{1}{2} \int_{\Omega \times [0, T]} \left(\mu'(t) - \frac{\mu(t)}{\tau} (|\nabla f(\mathbf{x})|^2 + C^2) \right) |\nabla h_t(\mathbf{x})|^2 d\mathbf{x} dt \\ & + \int_{\Omega \times [0, T]} \mu(t) \left(\hat{\eta}(\mathbf{x}) - \frac{\tau}{2} \right) (\Delta h_t(\mathbf{x}))^2 d\mathbf{x} dt \leq \frac{\tau \mu_{\max}}{2C^2} \int_{\Omega \times [0, T]} |\nabla \theta_t(\mathbf{x})|^2 d\mathbf{x} dt, \end{aligned} \quad (\text{A.13})$$

where $\mu_{\max} = \sup_{t \in [0, T]} \mu(t)$. We find a smooth function $\mu(t)$ such that

1. $\mu(t) \geq \mu_{\min} > 0$, for $t \in [0, T]$,
2. $\mu'(t) - \frac{2C^2}{\tau} \mu(t) \geq 0$.

A particular choice is

$$\mu(t) = \mu_{\min} e^{\frac{2C^2}{\tau} t}.$$

We then obtain from (A.13) that

$$\int_{\Omega \times [0, T]} \hat{\eta}(\mathbf{x}) (\Delta h_t(\mathbf{x}))^2 d\mathbf{x} dt \leq \frac{\tau \mu_{\max}}{\mu_{\min} C^2} \int_{\Omega \times [0, T]} |\nabla \theta_t(\mathbf{x})|^2 d\mathbf{x} dt. \quad (\text{A.14})$$

Now by employing (A.14) in bound (A.8) combined with (A.7) we get

$$\begin{aligned} \int_0^T \int_{\Omega} \theta_t(\mathbf{x}) u_t(\mathbf{x}) d\mathbf{x} dt & \leq \frac{1}{C} \sqrt{\frac{\tau \mu_{\max}}{\mu_{\min}}} \|\nabla \theta\|_{\mathcal{L}^2(\Omega \times [0, T])} \\ & \left(\int_{\Omega \times [0, T]} \frac{|\eta_t(\mathbf{x}) - \hat{\eta}_t(\mathbf{x})|^2}{\hat{\eta}_t(\mathbf{x})} u_t^2(\mathbf{x}) d\mathbf{x} dt \right)^{1/2}. \end{aligned} \quad (\text{A.15})$$

Next we note that

$$\begin{aligned} & \int_{\Omega \times [0, T]} |\eta_t(\mathbf{x}) - \hat{\eta}_t(\mathbf{x})|^2 u_t^2(\mathbf{x}) d\mathbf{x} dt \\ & \leq 2 \int_{\Omega \times [0, T]} |\eta_t^M(\mathbf{x}) - \hat{\eta}_t(\mathbf{x})|^2 u_t^2(\mathbf{x}) d\mathbf{x} dt \\ & + 2 \int_{\Omega \times [0, T]} ((\eta_t(\mathbf{x}) - M)_+)^2 u_t^2(\mathbf{x}) d\mathbf{x} dt. \end{aligned}$$

Call the first integral I_1 and denote the second one by I_2 . The integrand in I_2 is pointwise bounded by

$$2\eta_t^2(\mathbf{x}) u_t^2(\mathbf{x}) \mathbb{I}(\eta_t(\mathbf{x}) > M) \leq 2(\Phi(\rho_t(\mathbf{x})) - \Phi(\tilde{\rho}_t(\mathbf{x})))^2 \mathbb{I}(\eta_t(\mathbf{x}) > M).$$

Since $\rho_t, \tilde{\rho}_t \in \mathcal{L}^4$, we have that $(\Phi(\rho_t) - \Phi(\tilde{\rho}_t))^2$ has bounded integral. Hence, we can choose M large enough such that I_2 is arbitrarily small. Moreover we can choose the smooth approximation $\hat{\eta}_t$ such that I_1 is also arbitrarily small. Putting everything together, we obtain that

$$\int_{\Omega \times [0, T]} |\eta_t(\mathbf{x}) - \hat{\eta}_t(\mathbf{x})|^2 u_t^2(\mathbf{x}) d\mathbf{x} dt \leq \varepsilon,$$

where ε is an arbitrary small fixed constant.

In addition, since $\hat{\eta}_t(\mathbf{x}) \geq \tau$, invoking (A.15) we have

$$\int_0^T \int_{\Omega} \theta_t(\mathbf{x}) u_t(\mathbf{x}) d\mathbf{x} dt \leq \frac{1}{C} \sqrt{\frac{\mu_{\max}}{\mu_{\min}}} \|\nabla \theta\|_{\mathcal{L}^2(\Omega \times [0, T])} \sqrt{\varepsilon}.$$

Since $\frac{\mu_{\max}}{\mu_{\min}} = e^{\frac{2C^2}{\tau} T} < \infty$ and θ are independent of ε , by choosing ε arbitrarily small, we conclude that

$$\int_0^T \int_{\Omega} \theta_t(\mathbf{x}) u_t(\mathbf{x}) d\mathbf{x} dt \leq 0.$$

Since $\theta_t(\mathbf{x}) \geq 0$ was an arbitrary smooth function supported on $\Omega \times [0, T]$, this implies that $u \leq 0$, almost everywhere. By repeating a similar argument, we get $u \geq 0$, almost everywhere. The result follows. $\square$

B General results on the PDE (3.9) ($\delta > 0$)

This appendix contains some basic results on the PDE (3.9). Although these facts are standard, we collect them here for the reader's convenience.

In fact, we will consider a more general PDE, which also includes as a special case the one studied in [MMN18]. We consider a compact convex domain D , with a non-empty interior. The general PDE is parametrized by two functions $V \in \mathcal{C}^2(D)$ and $U \in \mathcal{C}^2(D \times D)$, with $U(\mathbf{x}_1, \mathbf{x}_2) = U(\mathbf{x}_2, \mathbf{x}_1)$. (Unlike in [MMN18], we consider the case of a compact domain with Neumann boundary conditions.) Given $\rho \in \mathcal{P}_2(D)$, we define

$$\Psi(\mathbf{x}, \rho) \equiv V(\mathbf{x}) + \int U(\mathbf{x}, \tilde{\mathbf{x}}) \rho(d\tilde{\mathbf{x}}), \quad (\text{B.1})$$

and consider the PDE

$$\partial_t \rho(\mathbf{x}, t) = \nabla \cdot (\rho(\mathbf{x}, t) \nabla \Psi(\mathbf{x}, \rho_t)) + \tau \Delta \rho(\mathbf{x}, t), \quad (\text{B.2})$$

with initial and boundary conditions

$$\begin{aligned} \rho_0 &= \rho_{\text{init}}, \\ \langle \mathbf{n}(\mathbf{x}), \rho_t(\mathbf{x}) \nabla \Psi(\mathbf{x}, \rho_t) + \tau \nabla \rho_t(\mathbf{x}) \rangle &= 0, \quad \forall \mathbf{x} \in \partial D. \end{aligned} \quad (\text{B.3})$$

We will typically write $\rho_t(\cdot)$ for a solution of this equation, in order to emphasize that it is a function of t that takes values in $\mathcal{P}_2(D)$, and $\rho(\mathbf{x}, t)$ for the corresponding density, viewed as a function on $D \times [0, T]$. Let us formally define the concept of weak solutions for the PDE (B.2).

Note that the PDE (3.9) is a special case of this setting with $D = \Omega^\delta$, and $V(\mathbf{w})$ and $U(\mathbf{w}_1, \mathbf{w}_2) = U(\mathbf{w}_1 - \mathbf{w}_2)$ defined as follows:

$$\begin{aligned} V(\mathbf{w}) &\equiv -\nu_0 K^\delta * f(\mathbf{w}) = -\nu_0 \int K^\delta(\mathbf{w} - \mathbf{x}) f(\mathbf{x}) d\mathbf{x}, \\ U(\mathbf{w}) &\equiv \nu_0 K^\delta * K^\delta(\mathbf{w}) = \nu_0 \int K^\delta(\mathbf{w} - \mathbf{x}) K^\delta(\mathbf{x}) d\mathbf{x}. \end{aligned} \tag{B.4}$$

Remark B.1. For the special choice of V and U given by (B.4) the following properties hold:

1. $V : \Omega^\delta \rightarrow \mathbb{R}$ is convex for any $\delta > 0$.
2. $\lim_{\delta \rightarrow 0} \sup_{\mathbf{w} \in \Omega^\delta} |V(\mathbf{w}) + \nu_0 f(\mathbf{w})| = 0$.
3. $U(\mathbf{w}) = \nu_0 \delta^{-2d} K^{(2)}(\mathbf{w}/\delta)$, where $K^{(2)} = K * K$.

Proof. We have $V(\mathbf{w}) = -\nu_0 \int K^\delta(\mathbf{x}) f(\mathbf{w} - \mathbf{x}) d\mathbf{x}$. Hence,

$$\begin{aligned} V(\lambda \mathbf{w} + (1 - \lambda) \mathbf{w}') &= -\nu_0 \int K^\delta(\mathbf{x}) f(\lambda \mathbf{w} + (1 - \lambda) \mathbf{w}' - \mathbf{x}) d\mathbf{x} \\ &= -\nu_0 \int K^\delta(\mathbf{x}) f(\lambda(\mathbf{w} - \mathbf{x}) + (1 - \lambda)(\mathbf{w}' - \mathbf{x})) d\mathbf{x} \\ &\leq -\nu_0 \int K^\delta(\mathbf{x}) (\lambda f(\mathbf{w} - \mathbf{x}) + (1 - \lambda) f(\mathbf{w}' - \mathbf{x})) d\mathbf{x} \\ &= \lambda V(\mathbf{w}) + (1 - \lambda) V(\mathbf{w}'). \end{aligned}$$

This proves that $V(\mathbf{w})$ is convex. The next two properties are straightforward. $\square$

Definition B.1 (Weak solution of PDE). *We say that $\rho : [0, T] \rightarrow \mathcal{P}_2(D)$ is a weak solution of (B.2) with initial and boundary conditions (B.3) if $\rho \in \mathcal{C}([0, T], \mathcal{P}_2(D))$ and, for any test function $h \in \mathcal{C}^{2,1}(D \times [0, T])$, satisfying $\langle \mathbf{n}(\mathbf{x}), \nabla h(\mathbf{x}, t) \rangle = 0$ for all $\mathbf{x} \in \partial D, t \in [0, T]$, we have*

$$\begin{aligned} &\int_D h(\mathbf{x}, T) \rho_T(d\mathbf{x}) - \int_D h(\mathbf{x}, 0) \rho_{\text{init}}(d\mathbf{x}) \\ &= \int_0^T \int_D [\partial_t h(\mathbf{x}, t) + \tau \Delta h(\mathbf{x}, t) - \langle \nabla \Psi(\mathbf{x}, \rho_t), \nabla h(\mathbf{x}, t) \rangle] \rho_t(d\mathbf{x}) dt. \end{aligned} \tag{B.5}$$

We now state and prove Duhamel's principle for the PDE (B.2). Duhamel's principle follows from the fact that the right-hand side of (B.2) contains the linear diffusion term $\tau \Delta \rho$, and it will be crucial for the proofs that will follow.

Lemma B.2 (Duhamel's principle). *Assume $\tau > 0$. Let $G^D(\mathbf{x}, \mathbf{y}; t)$ denote the heat kernel with Neumann boundary conditions, defined in (G.1)-(G.3). Let ρ be a weak solution of the PDE (B.2) with initial and boundary conditions (B.3). Then, for any $t > 0$, $\rho_t(d\mathbf{x})$ has a density, denoted by $\rho(\cdot, t)$, which satisfies, for any $t > 0$,*

$$\begin{aligned} \rho(\mathbf{x}, t) &= \int_D G^D(\mathbf{x}, \mathbf{y}; \tau t) \rho_{\text{init}}(d\mathbf{y}) \\ &\quad - \int_0^t \int_D \langle \nabla_{\mathbf{y}} G^D(\mathbf{x}, \mathbf{y}; \tau(t - s)), \nabla_{\mathbf{y}} \Psi(\mathbf{y}; \rho_s) \rangle \rho(\mathbf{y}, s) d\mathbf{y} ds. \end{aligned} \tag{B.6}$$

Proof. By rescaling time, without loss of generality, we set $\tau = 1$. Let $\varphi \in \mathcal{C}^2(D)$, and define

$$G_\varphi(\mathbf{x}; t) = \int_D G^D(\mathbf{x}, \mathbf{y}; t) \varphi(\mathbf{y}) d\mathbf{y}. \quad (\text{B.7})$$

By the properties of the heat kernel, we have:

$$(\partial_t - \Delta) G_\varphi(\mathbf{x}; t) = 0 \quad \forall t > 0, \quad (\text{B.8})$$

$$\langle \mathbf{n}(\mathbf{x}), \nabla G_\varphi(\mathbf{x}; t) \rangle = 0 \quad \forall \mathbf{x} \in \partial D, \quad (\text{B.9})$$

$$\lim_{t \rightarrow 0} G_\varphi(\mathbf{x}; t) = G_\varphi(\mathbf{x}; 0) = \varphi(\mathbf{x}). \quad (\text{B.10})$$

Let ρ_t be a weak solution. We choose the test function $h(\mathbf{x}, s) = G_\varphi(\mathbf{x}; t - s)$ in (B.5) with $T = t$. Note that by (B.8), this test function satisfies the Neumann boundary condition. In addition, by (B.9) we obtain

$$\begin{aligned} \int_D \varphi(\mathbf{y}) \rho_t(d\mathbf{y}) &= \int_D G_\varphi(\mathbf{x}; t) \rho_{\text{init}}(d\mathbf{x}) \\ &\quad - \int_0^t \int_D \langle \nabla \Psi(\mathbf{x}, \rho_s), \nabla G_\varphi(\mathbf{x}; t - s) \rangle \rho_s(d\mathbf{x}) ds. \end{aligned} \quad (\text{B.11})$$

By an application of Fubini's theorem, this implies

$$\begin{aligned} \int_D \varphi(\mathbf{y}) \rho_t(d\mathbf{y}) &= \int_D \int_D G^D(\mathbf{x}, \mathbf{y}; t) \rho_{\text{init}}(d\mathbf{x}) \varphi(\mathbf{y}) d\mathbf{y} \\ &\quad - \int_0^t \int_D \int_D \langle \nabla \Psi(\mathbf{x}, \rho_s), \nabla G^D(\mathbf{x}, \mathbf{y}; t - s) \rangle \varphi(\mathbf{y}) \rho_s(d\mathbf{x}) ds d\mathbf{y}. \end{aligned} \quad (\text{B.12})$$

Since $\varphi \in \mathcal{C}^2(D)$ is arbitrary, we obtain that ρ_t admits a density and (B.6) follows. $\square$

As an intermediate step towards proving existence and uniqueness, we consider a linearized problem

$$\partial_t \rho(\mathbf{x}, t) = \nabla \cdot (\rho(\mathbf{x}, t) \nabla \Psi_*(\mathbf{x}, t)) + \tau \Delta \rho(\mathbf{x}, t), \quad (\text{B.13})$$

with initial and boundary conditions

$$\begin{aligned} \rho_0 &= \rho_{\text{init}}, \\ \langle \mathbf{n}(\mathbf{x}), \rho_t(\mathbf{x}) \nabla \Psi_*(\mathbf{x}, t) + \tau \nabla \rho_t(\mathbf{x}) \rangle &= 0, \quad \forall \mathbf{x} \in \partial D. \end{aligned} \quad (\text{B.14})$$

Here, $\Psi_* : D \times \mathbb{R} \rightarrow \mathbb{R}$ is independent of ρ , and weak solutions are defined as for the original problem (with Neumann boundary conditions).

Corollary B.3 (Uniqueness of linearized problem). *Assume that $\tau > 0$ and also that*

$$\|\nabla \Psi_*\|_{\mathcal{L}^\infty(D \times [0, T])} \equiv \sup_{t \in [0, T]} \sup_{\mathbf{x} \in D} |\nabla \Psi_*(\mathbf{x}, t)| < \infty.$$

Then, the PDE (B.13) with initial and boundary conditions (B.14) has at most one weak solution.

Proof. Without loss of generality, we will set $\tau = 1$. Assume by contradiction that $\rho^{(1)}, \rho^{(2)}$ are two solutions. Fix arbitrary $0 \leq t' \leq t$. Then, by an application of (B.6) to $\Psi_*(\mathbf{x}, t)$, we have

$$\begin{aligned}
& |\rho^{(1)}(\mathbf{x}, t') - \rho^{(2)}(\mathbf{x}, t')| \leq \|\rho^{(1)}(\cdot, 0) - \rho^{(2)}(\cdot, 0)\|_{\mathcal{L}^\infty(D)} \\
& + \|\nabla \Psi_*\|_{\mathcal{L}^\infty(D \times [0, T])} \int_0^{t'} \int_D |\nabla_{\mathbf{y}} G^D(\mathbf{x}, \mathbf{y}; t' - s)| |\rho^{(1)}(\mathbf{y}, s) - \rho^{(2)}(\mathbf{y}, s)| d\mathbf{y} ds \\
& \leq \|\rho^{(1)}(\cdot, 0) - \rho^{(2)}(\cdot, 0)\|_{\mathcal{L}^\infty(D)} \\
& + C(D) \|\nabla \Psi_*\|_{\mathcal{L}^\infty(D \times [0, T])} \int_0^{t'} \frac{1}{\sqrt{t' - s}} \|\rho^{(1)}(\cdot, s) - \rho^{(2)}(\cdot, s)\|_{\mathcal{L}^\infty(D)} ds \\
& \leq \|\rho^{(1)}(\cdot, 0) - \rho^{(2)}(\cdot, 0)\|_{\mathcal{L}^\infty(D)} \\
& + C(D) \|\nabla \Psi_*\|_{\mathcal{L}^\infty(D \times [0, T])} \sqrt{t} \sup_{s \leq t} \|\rho^{(1)}(\cdot, s) - \rho^{(2)}(\cdot, s)\|_{\mathcal{L}^\infty(D)},
\end{aligned}$$

where we used the estimates of Theorem G.1. By taking supremum over $0 \leq t' \leq t$ from both sides, we obtain that for $t < 1/(C(D)^2 \|\nabla \Psi_*\|_{\mathcal{L}^\infty(D \times [0, T])}^2)$,

$$\sup_{s \leq t} \|\rho^{(1)}(\cdot, s) - \rho^{(2)}(\cdot, s)\|_{\mathcal{L}^\infty(D)} \leq \frac{\|\rho^{(1)}(\cdot, 0) - \rho^{(2)}(\cdot, 0)\|_{\mathcal{L}^\infty(D)}}{1 - C(D) \|\nabla \Psi_*\|_{\mathcal{L}^\infty(D \times [0, T])} \sqrt{t}}. \quad (\text{B.15})$$

Therefore, the two solutions coincide if we fix the initial condition $\rho^{(1)}(\cdot, 0) = \rho^{(2)}(\cdot, 0) = \rho_{\text{init}}$. For larger t , the claim follows by iterating the above argument. $\square$

C Nonlinear dynamics

The ‘nonlinear dynamics’ plays an important role in our proof of Theorem 5.1. In this section we adopt the same general setting as in Appendix B, remembering that for our application we set $D = \Omega^\delta$ and U, V as per Eq. (B.4).

Given $\rho : [0, T] \rightarrow \mathcal{P}_2(D)$, consider the following stochastic differential equation for a process $(\mathbf{X}_t)_{t \in [0, T]}$, with a reflecting boundary condition (known as ‘Skorokhod problem’)

$$\mathbf{X}_0 \sim \rho_{\text{init}} \quad (\text{C.1})$$

$$d\mathbf{X}_t = -\nabla \Psi(\mathbf{X}_t, \rho_t) dt + \sqrt{2\tau} d\mathbf{B}_t + d\Phi_t, \quad (\text{C.2})$$

where $(\mathbf{B}_t)_{t \geq 0}$ is a standard d -dimensional Brownian motion and $(\Phi_t)_{t \geq 0}$ enforces the reflecting boundary by satisfying the following constraints (recall that $\mathbf{n}(\mathbf{x})$ is the normal to ∂D at $\mathbf{x} \in \partial D$, directed inside):

- (i) $(\Phi_t)_{t \geq 0}$ is adapted (and hence so is $(\mathbf{X}_t)_{t \geq 0}$).
- (ii) $t \mapsto \Phi_t$ has (almost surely) bounded variation. Denoting by $\|\Phi\|_{\text{TV}}(t)$ the total variation of Φ on the interval $[0, t]$, we define the measure μ_Φ on $[0, T]$ by $\mu_\Phi([0, t]) = \|\Phi\|_{\text{TV}}(t)$.
- (iii) $\mu_\Phi(\{t : \mathbf{X}_t \in D^\circ\}) = 0$, where D° denotes the interior of D .

(iv) We have that, for $t \in [0, T]$,

$$\Phi_t = \int_0^t N_s \mu_\Phi(ds), \quad (\text{C.3})$$

where $N_s = \mathbf{n}(\mathbf{X}_s)$, for μ_Φ -almost every s .

Then, $(\mathbf{X}_t, \Phi_t)_{t \in [0, T]}$ is said to solve the Skorokhod problem.

Lemma C.1 (Existence, uniqueness and continuity of Skorokhod problem). *Fix $\rho_{\text{init}} \in \mathcal{P}_2(D)$ and let $\rho : [0, T] \rightarrow \mathcal{P}_2(D)$ with $\rho_0 = \rho_{\text{init}}$. Then, the Skorokhod problem (C.1), (C.2) admits a unique solution $(\mathbf{X}_t)_{t \geq 0}$ with continuous paths. Define $\mathcal{F}(\rho)_t \in \mathcal{P}_2(D)$, for $t \in [0, T]$, by letting $\mathcal{F}(\rho)_t = \text{Law}(\mathbf{X}_t)$. Then, $\mathcal{F}(\rho) \in \mathcal{C}([0, T], \mathcal{P}_2(D))$.*

Proof. Let $\mathbf{b}(\mathbf{x}, t) \equiv -\nabla \Psi(\mathbf{x}, \rho_t)$ and notice that, by the smoothness of U, V , and compactness of D , this is a Lipschitz continuous function of $\mathbf{x}$. Hence the problem (C.1), (C.2) admits a unique solution by [Tan79, Theorem 4.1].

We are left with the task of proving that $t \mapsto \mathcal{F}(\rho)_t$ is continuous in W_2 metric. Notice that

$$\mathbf{X}_t = \mathbf{X}_0 + \int_0^t \mathbf{b}(\mathbf{X}_s, s) ds + \sqrt{2\tau} \mathbf{B}_t + \Phi_t. \quad (\text{C.4})$$

By [Tan79, Lemma 2.2], we have, for any $s \leq t$,

$$\begin{aligned} |\mathbf{X}_t - \mathbf{X}_s|^2 &\leq \left| \int_s^t \mathbf{b}(\mathbf{X}_r, r) dr + \sqrt{2\tau} (\mathbf{B}_t - \mathbf{B}_s) \right|^2 + \\ &\quad 2 \int_s^t \left\langle \int_r^t \mathbf{b}(\mathbf{X}_u, u) du + \sqrt{2\tau} (\mathbf{B}_t - \mathbf{B}_r), \mathbf{N}_r \right\rangle \mu_\Phi(dr). \end{aligned}$$

Taking expectation, we get

$$\begin{aligned} \mathbb{E}\{|\mathbf{X}_t - \mathbf{X}_s|^2\} &\leq \sup_{\mathbf{x}, t} |\mathbf{b}(\mathbf{x}, t)|^2 (t-s)^2 + 2\tau(t-s) \\ &\quad + 2 \sup_{\mathbf{x}, t} |\mathbf{b}(\mathbf{x}, t)| \int_s^t (t-r) \mu_\Phi(dr) \\ &\leq \sup_{\mathbf{x}, t} |\mathbf{b}(\mathbf{x}, t)|^2 (t-s)^2 + 2\tau(t-s) \\ &\quad + 2 \sup_{\mathbf{x}, t} |\mathbf{b}(\mathbf{x}, t)| (t-s) \mu_\Phi([0, t]), \end{aligned} \quad (\text{C.5})$$

whence the continuity follows. $\square$

Definition C.2 (Solution of nonlinear dynamics). *We say that $\rho \in \mathcal{C}([0, T]; \mathcal{P}_2(D))$ is a solution of the nonlinear dynamics if $\mathcal{F}(\rho) = \rho$, namely*

$$\rho_t = \text{Law}(\mathbf{X}_t) \quad \forall t \in [0, T]. \quad (\text{C.6})$$

Lemma C.3. *Assume $\tau > 0$. If $\rho : [0, T] \rightarrow \mathcal{P}_2(D)$ is a weak solution of the PDE (B.2) with initial and boundary conditions (B.3), then it is a solution of the nonlinear dynamics. Vice versa, if $\rho : [0, T] \rightarrow \mathcal{P}_2(D)$ is a solution of the nonlinear dynamics, then it is a weak solution of PDE (B.2) with initial and boundary conditions (B.3).*

Proof. Let ρ be a weak solution of the PDE (B.2), and assume $\tau > 0$. Let $(\mathbf{X}_t)_{t \geq 0}$ be the unique solution of the Skorokhod problem (C.1), (C.2), cf. Lemma C.1. Let $\tilde{\rho}_t \equiv \text{Law}(\mathbf{X}_t)$, $t \geq 0$, i.e. $\tilde{\rho} \equiv \mathcal{F}(\rho)$. For $g \in \mathcal{C}^2(D)$, satisfying $\langle \mathbf{n}(\mathbf{x}), \nabla g(\mathbf{x}) \rangle = 0$ for all $\mathbf{x} \in \partial D$, compute

$$\begin{aligned}
\int g(\mathbf{x}) \tilde{\rho}_t(d\mathbf{x}) &= \mathbb{E}\{g(\mathbf{X}_t)\} \\
&\stackrel{(a)}{=} \mathbb{E}\{g(\mathbf{X}_0)\} + \int_0^t \mathbb{E}\{-\langle \nabla \Psi(\mathbf{X}_s, \rho_s), \nabla g(\mathbf{X}_s) \rangle + \tau \Delta g(\mathbf{X}_s)\} ds \\
&\quad + \mathbb{E} \int_0^t \langle \nabla g(\mathbf{X}_s), \mathbf{N}_s \rangle \mu_\Phi(ds) \\
&\stackrel{(b)}{=} \mathbb{E}\{g(\mathbf{X}_0)\} + \int_0^t \mathbb{E}\{-\langle \nabla \Psi(\mathbf{X}_s, \rho_s), \nabla g(\mathbf{X}_s) \rangle + \tau \Delta g(\mathbf{X}_s)\} ds \\
&\stackrel{(c)}{=} \int_D g(\mathbf{x}) \rho_{\text{init}}(d\mathbf{x}) \\
&\quad + \int_0^t \int_D \{-\langle \nabla \Psi(\mathbf{x}, \rho_s), \nabla g(\mathbf{x}) \rangle + \tau \Delta g(\mathbf{x})\} \tilde{\rho}_s(d\mathbf{x}) ds.
\end{aligned} \tag{C.7}$$

Here (a) follows from Ito's formula for continuous semimartingales [RW94], (b) since $\mathbf{X}_s \in \partial D$ and $\mathbf{N}_s = \mathbf{n}(\mathbf{X}_s)$ for μ_Φ -almost every s , and (c) by the definition of $\tilde{\rho}$. We conclude that $\tilde{\rho}$ is a weak solution of the linearized PDE (B.13), with $\Psi_*(\mathbf{x}, t) = \Psi(\mathbf{x}, \rho_t)$. Since ρ also solves the same linearized PDE, we conclude by Lemma B.3 that $\tilde{\rho}_t = \rho_t$ for all $t \in [0, T]$, and therefore ρ is a solution of the nonlinear dynamics.

Next, assume that $\rho : [0, T] \rightarrow \mathcal{P}_2(D)$ is a solution of the nonlinear dynamics. Then by the same application of Ito's formula to the process $\mathbf{X}_t$, we have

$$\begin{aligned}
\int g(\mathbf{x}) \rho_t(d\mathbf{x}) &= \mathbb{E}\{g(\mathbf{X}_0)\} \\
&\quad + \int_0^t \mathbb{E}\{-\langle \nabla \Psi(\mathbf{X}_s, \rho_s), \nabla g(\mathbf{X}_s) \rangle + \tau \Delta g(\mathbf{X}_s)\} ds,
\end{aligned} \tag{C.8}$$

which coincides with the claim that ρ is a weak solution of the PDE (B.2). □

Theorem C.4 (Existence and uniqueness of nonlinear dynamics). *For any initial condition $\rho_{\text{init}} \in \mathcal{P}_2(D)$, and any $T > 0$, the nonlinear dynamics (C.6) admits a unique solution $\rho : [0, T] \rightarrow \mathcal{P}_2(D)$ with $\rho_0 = \rho_{\text{init}}$. As a consequence, the PDE (B.2) with initial and boundary conditions (B.3) has a unique solution.*

Proof. Note that it is sufficient to prove the claim for $T \leq T_0$, where $T_0 > 0$ is a small enough constant, since this implies the claim for arbitrary T by breaking $[0, T]$ into intervals of size smaller than T_0 .

We claim that $\mathcal{F}$ is a contraction on $\mathcal{C}([0, T], \mathcal{P}_2(D))$ endowed with the metric $d(\rho, \tilde{\rho}) \equiv \sup_{t \in [0, T]} W_2(\rho_t, \tilde{\rho}_t)$. To show that this is the case, define $\mathbf{b}(\mathbf{x}, t) \equiv -\nabla \Psi(\mathbf{x}, \rho_t)$, $\tilde{\mathbf{b}}(\mathbf{x}, t) \equiv -\nabla \Psi(\mathbf{x}, \tilde{\rho}_t)$. By the smoothness of U , V and by the compactness of D , we have that $\mathbf{b}$ and $\tilde{\mathbf{b}}$ are Lipschitz continuous in $\mathbf{x}$, with Lipschitz constant L independent of $t, \rho, \tilde{\rho}$. Further,

$$\begin{aligned}
|\mathbf{b}(\mathbf{x}, t) - \tilde{\mathbf{b}}(\mathbf{x}, t)| &= \left| \int \nabla U(\mathbf{x}, \tilde{\mathbf{x}}) \rho_t(d\tilde{\mathbf{x}}) - \int \nabla U(\mathbf{x}, \tilde{\mathbf{x}}) \tilde{\rho}_t(d\tilde{\mathbf{x}}) \right| \\
&\leq C W_1(\rho_t, \tilde{\rho}_t) \leq C d(\rho, \tilde{\rho}).
\end{aligned} \tag{C.9}$$

Let $(\mathbf{X}_t, \Phi_t)$ and $(\tilde{\mathbf{X}}, \tilde{\Phi}_t)$ are be solution of the Skorokhod problem (C.2), with drift coefficients $\mathbf{b}(\mathbf{x}, t)$, $\tilde{\mathbf{b}}(\mathbf{x}, t)$. We couple the processes $\mathbf{X}_t$ and $\tilde{\mathbf{X}}_t$ by using the same initial condition $\mathbf{X}_0$ and same Brownian motion $\mathbf{B}_t$:

$$\mathbf{X}_t = \mathbf{X}_0 + \int_0^t \mathbf{b}(\mathbf{X}_s, s) ds + \sqrt{2\tau} \mathbf{B}_t + \Phi_t, \quad (\text{C.10})$$

$$\tilde{\mathbf{X}}_t = \mathbf{X}_0 + \int_0^t \tilde{\mathbf{b}}(\tilde{\mathbf{X}}_s, s) ds + \sqrt{2\tau} \mathbf{B}_t + \tilde{\Phi}_t. \quad (\text{C.11})$$

Define

$$\mathbf{D}_t \equiv \int_0^t \mathbf{b}(\mathbf{X}_s, s) ds - \int_0^t \tilde{\mathbf{b}}(\tilde{\mathbf{X}}_s, s) ds, \quad (\text{C.12})$$

and notice that, by the above remarks,

$$|\mathbf{D}_t| \leq L \int_0^t |\mathbf{X}_s - \tilde{\mathbf{X}}_s| ds + C t d(\rho, \tilde{\rho}). \quad (\text{C.13})$$

Further, by [Tan79, Remark 2.2], we have

$$\begin{aligned} |\mathbf{X}_t - \tilde{\mathbf{X}}_t|^2 &\leq 2 \int_0^t \langle \mathbf{X}_s - \tilde{\mathbf{X}}_s, \mathbf{D}_s \rangle ds \\ &\leq 2 \int_0^t |\mathbf{X}_s - \tilde{\mathbf{X}}_s| \left(L \left(\int_0^s |\mathbf{X}_r - \tilde{\mathbf{X}}_r| dr \right) + C s d(\rho, \tilde{\rho}) \right) ds \\ &\leq 2L \left(\int_0^t |\mathbf{X}_s - \tilde{\mathbf{X}}_s| ds \right)^2 + 2C t d(\rho, \tilde{\rho}) \int_0^t |\mathbf{X}_s - \tilde{\mathbf{X}}_s| ds \\ &\leq 2Lt \int_0^t |\mathbf{X}_s - \tilde{\mathbf{X}}_s|^2 ds + 2C t^{3/2} d(\rho, \tilde{\rho}) \left(\int_0^t |\mathbf{X}_s - \tilde{\mathbf{X}}_s|^2 ds \right)^{1/2}. \end{aligned} \quad (\text{C.14})$$

Define $\Delta(t) \equiv \mathbb{E}\{|\mathbf{X}_t - \tilde{\mathbf{X}}_t|^2\}$ and $\bar{\Delta}(t) \equiv \sup_{s \leq t} \Delta(s)$. By taking the expectation of the last inequality and using Jensen's inequality, we get

$$\Delta(t) \leq 2Lt \int_0^t \Delta(s) ds + 2C t^{3/2} d(\rho, \tilde{\rho}) \left(\int_0^t \Delta(s) ds \right)^{1/2}, \quad (\text{C.15})$$

which immediately implies

$$\bar{\Delta}(t) \leq 2Lt^2 \bar{\Delta}(t) + 2C t^2 d(\rho, \tilde{\rho}) \bar{\Delta}(t)^{1/2}. \quad (\text{C.16})$$

Hence, for $T_0 < (2L)^{-1/2}$,

$$\bar{\Delta}(t) \leq \left(\frac{2CT_0^2}{1 - 2LT_0^2} \right)^2 d(\rho, \tilde{\rho})^2. \quad (\text{C.17})$$

Selecting T_0 small enough, so that $(2CT_0^2)/(1 - 2LT_0^2) \leq 1/2$, we obtain

$$d(\mathcal{F}(\rho), \mathcal{F}(\tilde{\rho})) \leq \sqrt{\bar{\Delta}(T_0)} \leq \frac{1}{2} d(\rho, \tilde{\rho}). \quad (\text{C.18})$$

This proves that $\mathcal{F}$ is a contraction as claimed. By Lemma C.1, $\mathcal{F}$ maps $\mathcal{C}([0, T], \mathcal{P}_2(D))$ into itself. Furthermore, $\mathcal{C}([0, T], \mathcal{P}_2(D))$ is complete with respect to the metric d . As a result, there exists a unique fixed point. $\square$

We conclude this section by stating a result about the discretization of the nonlinear dynamics. Fix a solution $(\rho_t)_{t \geq 0}$ of the PDE (B.2) with initial condition $\rho_0 = \rho_{\text{init}}$, a step size $\varepsilon > 0$ and define recursively the random variables $(\mathbf{X}^\varepsilon)_{k \in \mathbb{N}}$ by

$$\mathbf{X}_0^\varepsilon \sim \rho_{\text{init}} \quad (\text{C.19})$$

$$\mathbf{X}_{k+1}^\varepsilon = \mathbf{P} \left(\mathbf{X}_k^\varepsilon - \varepsilon \nabla \Psi(\mathbf{X}_k^\varepsilon, \rho_{k\varepsilon}) + \sqrt{2\tau\varepsilon} \mathbf{g}_k \right). \quad (\text{C.20})$$

This can be viewed as an Euler discretization of the stochastic differential equation (C.1), (C.2), and the next theorem establishes that this is indeed a close approximation of the original process. It is just an immediate consequence of a result of Słomiński [Slo94, Slo01].

Theorem C.5 (Theorem 3.2 in [Slo01]). *Consider the nonlinear dynamics defined by Eqs. (C.1), (C.2). Assume $\mathbf{B}(\mathbf{0}; r) \subseteq D$, and $\|\nabla V\|_{\mathcal{L}^\infty(D)}$, $\|\nabla U\|_{\mathcal{L}^\infty(D \times D)}$, $\|\nabla V\|_{\text{Lip}}$, $\|\nabla U\|_{\text{Lip}} \leq L$. Also assume that $\text{supp}(\rho_{\text{init}}) \subseteq \mathbf{B}(\mathbf{0}, r)$. Construct the Euler scheme (C.19), (C.20) on the same probability space by letting $\mathbf{X}_0^\varepsilon = \mathbf{X}_0$ and $\mathbf{g}_k = (\mathbf{B}((k+1)\varepsilon) - \mathbf{B}(k\varepsilon))/\sqrt{\varepsilon}$. Then, for any $p \in \mathbb{N}$, $T \in \mathbb{R}_{\geq 0}$,*

$$\mathbb{E} \left\{ \max_{k \in [0, T/\varepsilon] \cap \mathbb{N}} |\mathbf{X}_k^\varepsilon - \mathbf{X}_{k\varepsilon}|^{2p} \right\}^{1/(2p)} \leq C_* L r e^{C_* p L T} (d^2 \varepsilon \log(1/\varepsilon))^{1/4}. \quad (\text{C.21})$$

Proof. The proof is obtained simply by chasing the constants in the proof of Theorem 3.2 (part (ii)) of [Slo01], and using the optimal constant in the Burkholder-Davis-Gundy inequality (which yields $C(p) \leq (C_* p)^{2p}$ in [Slo01, Eq. (2.7)]). $\square$

D Convergence of SGD to the PDE: Proof of Theorem 5.1

The proof is a ‘propagation of chaos’ argument [Szn91]. While the basic idea is similar to the one used in [MMN18], implementing it requires different estimates because of the reflecting boundary conditions. In particular, we rely on tools developed in the study of discretizations of reflecting stochastic differential equations.

We will prove a more general theorem that implies Theorem 5.1 as a special case, and also applies to the setting of [MMN18]. Namely, we consider data $\{\mathbf{z}_i = (y_i, \mathbf{x}_i)\}_{i \geq 1}$ i.i.d. with common distribution $\mathbb{P}$ on $\mathbb{R} \times \mathbb{R}^{d_0}$, and parameters $\mathbf{w}_i \in D \subseteq \mathbb{R}^d$. These parameters are initially sampled independently from distribution $\rho_0 \in \mathcal{P}_2(D)$, and then evolve according to

$$\mathbf{w}_i^{k+1} = \mathbf{P} \left\{ \mathbf{w}_i^k + \mathbf{F}_i(\mathbf{z}_{k+1}; \mathbf{w}^k) \right\}, \quad (\text{D.1})$$

$$\mathbf{F}_i(\mathbf{z}_{k+1}; \mathbf{w}^k) = -\varepsilon \nabla \sigma(\mathbf{x}_{k+1}; \mathbf{w}_i^k) \left(y_{k+1} - \frac{1}{N} \sum_{i=1}^N \sigma(\mathbf{x}_{k+1}; \mathbf{w}_i^k) \right) + \sqrt{2\tau\varepsilon} \mathbf{g}_i^{k+1}. \quad (\text{D.2})$$

Here $\mathbf{P}$ is the projection on the closed convex domain $D \subseteq \mathbb{R}^d$ with non-empty interior. The setting of Theorem 5.1 is recovered by taking $\sigma(\mathbf{x}; \mathbf{w}) = K_\delta(\mathbf{x} - \mathbf{w})$, $D = \Omega^\delta$, $\mathbf{x}_k \sim \text{Unif}(\Omega)$, $\mathbb{E}\{y_k | \mathbf{x}_k\} = f(\mathbf{x}_k)$.

We make the following assumptions:

(G1) $\|y\|_\infty$, $\|\sigma\|_\infty = \text{ess sup}_{\mathbf{w} \in D, \mathbf{x}} |\sigma(\mathbf{x}; \mathbf{w})| \leq \sigma_\infty$, and $\nabla_{\mathbf{w}} \sigma(\mathbf{x}; \mathbf{w})$ is γ -subgaussian.

(G2) Letting $V(\mathbf{w}) = -\mathbb{E}\{y\sigma(\mathbf{x}; \mathbf{w})\}$, $U(\mathbf{w}_1, \mathbf{w}_2) \equiv \mathbb{E}\{\sigma(\mathbf{x}; \mathbf{w}_1)\sigma(\mathbf{x}; \mathbf{w}_2)\}$, both V and U are differentiable with Lipschitz continuous derivative, namely $\|\nabla V\|_{\text{Lip}}, \|\nabla U\|_{\text{Lip}} \leq L$. Further, we assume $\|\nabla U\|_{\mathcal{L}^\infty(D \times D)} < \infty$.

Theorem D.1. *Consider the general update (D.1) with initialization $(\mathbf{w}_i^0)_{i \leq N} \sim_{\text{iid}} \rho_0 = \rho_{\text{init}}$, under the conditions (G1), (G2) above. For $t \geq 0$, let ρ_t be the unique solution of the PDE (B.2) with initial and boundary conditions (B.3). Assume $\text{supp}(\rho_{\text{init}}) \subseteq \mathbf{B}(\mathbf{0}, r)$.*

Then, for $T \geq 0$ $TL \geq 1$, any $g : \mathbb{R}^d \rightarrow \mathbb{R}$ with $\|g\|_{\text{Lip}} \leq 1$ and for $\varepsilon \leq 1$, $p \in \mathbb{N}$, the following holds with probability at least $1 - z^{-2p}$:

$$\begin{aligned} \sup_{k \in [0, T/\varepsilon] \cap \mathbb{N}} \left| \sum_{i=1}^N g(\mathbf{w}_i^k) - \int g(\mathbf{w}) \rho_{k\varepsilon}(\mathrm{d}\mathbf{w}) \right| &\leq z \text{err}(N, d, \varepsilon) e^{C_* p L T}, \\ \sup_{k \in [0, T/\varepsilon] \cap \mathbb{N}} |R_N(\mathbf{w}^k) - R^\delta(\rho_{k\varepsilon})| &\leq z \text{err}(N, d, \varepsilon) e^{C_* p L T}, \end{aligned} \quad (\text{D.3})$$

where

$$\text{err}(N, d, \varepsilon) = \sqrt{\frac{d}{N}} \vee (\sigma_\infty \gamma \sqrt{d\varepsilon}) \vee \left(Lr(d^2 \varepsilon \log(1/\varepsilon))^{1/4} \right). \quad (\text{D.4})$$

Theorem 5.1 follows as a special case of Theorem D.1 by considering $\sigma(\mathbf{x}; \mathbf{w}) = K_\delta(\mathbf{x} - \mathbf{w})$ and letting $\sigma_\infty \leq C_* \delta^{-d}$, $\gamma = C_* \delta^{-d-1}$ and $L = C_* \delta^{-2d-1}$.

Proof. Let $\mathcal{F}_k$ denote the sigma algebra generated by $(\mathbf{z}_j)_{j \leq k}$ and denote the empirical distribution of $(\mathbf{w}_i^k)_{i \leq N}$ by $\rho_k^{(N)} \equiv \sum_{i=1}^N \delta_{\mathbf{w}_i^k}$. Note that

$$\begin{aligned} \mathbb{E}\{\mathbf{F}_i(\mathbf{z}_{k+1}; \mathbf{w}^k) | \mathcal{F}_k\} &= \varepsilon \mathbf{G}(\mathbf{w}_i^k; \rho_k^{(N)}), \\ \mathbf{G}(\mathbf{w}; \rho) &\equiv -\nabla \Psi(\mathbf{w}; \rho) = -\nabla V(\mathbf{w}) - \int \nabla U(\mathbf{w}, \mathbf{w}') \rho(\mathrm{d}\mathbf{w}'). \end{aligned} \quad (\text{D.5})$$

We introduce two auxiliary processes $(\bar{\mathbf{w}}_i^k)_{i \leq N}$ $(\hat{\mathbf{w}}_i^k)_{i \leq N}$, with initial conditions $\bar{\mathbf{w}}_i^0 = \hat{\mathbf{w}}_i^0 = \mathbf{w}_i^0$, as follows:

- The trajectories $(\hat{\mathbf{w}}_i^k)_{k \geq 0}$ are i.i.d. copies of the nonlinear dynamics introduced in Appendix C, sampled at times $t = k\varepsilon$. Namely, for any $k \in \mathbb{R}$

$$\hat{\mathbf{w}}_i^k = \mathbf{w}_i^0 + \int_0^{k\varepsilon} \mathbf{G}(\mathbf{w}_i^{s/\varepsilon}; \rho_s) \mathrm{d}s + \sqrt{2\tau} \mathbf{B}_i(k\varepsilon) + \Phi_i(k\varepsilon). \quad (\text{D.6})$$

In particular, for any k , $(\hat{\mathbf{w}}_i^k)_{i \leq N} \sim_{\text{iid}} \rho_{k\varepsilon}$.

- The trajectories $(\bar{\mathbf{w}}_i^k)_{k \geq 0}$ are obtained by the Euler discretization of the non-linear dynamics:

$$\bar{\mathbf{w}}_i^{k+1} = \mathbf{P}\left(\bar{\mathbf{w}}_i^k + \varepsilon \mathbf{G}(\bar{\mathbf{w}}_i^k; \rho_{k\varepsilon}) + \sqrt{2\tau \varepsilon} \mathbf{g}_i^{k+1}\right). \quad (\text{D.7})$$

As above, $(\rho_s)_{s \geq 0}$ is the solution of the PDE (B.2). Note that, again, the $(\bar{\mathbf{w}}_i^k)_{i \leq N}$ are i.i.d. although their distribution does not coincide with $\rho_{k\varepsilon}$.

We construct these three processes on the same space by letting $\mathbf{B}_i((k+1)\varepsilon) = \mathbf{B}_i(k\varepsilon) + \sqrt{\varepsilon}\mathbf{g}_i^{k+1}$, and define the distances (for $q \geq 1$)

$$\mathcal{D}_q(k) \equiv \mathbb{E} \left\{ \max_{j \leq k} |\mathbf{w}_i^j - \bar{\mathbf{w}}_i^j|^q \right\}^{1/q} = \mathbb{E} \left\{ \frac{1}{N} \sum_{i=1}^N \max_{j \leq k} |\mathbf{w}_i^j - \bar{\mathbf{w}}_i^j|^q \right\}^{1/q} \quad (\text{D.8})$$

$$\geq \mathbb{E} \left\{ \max_{j \leq k} \frac{1}{N} \sum_{i=1}^N |\mathbf{w}_i^j - \bar{\mathbf{w}}_i^j|^q \right\}^{1/q},$$

$$\widehat{\mathcal{D}}_q(k) \equiv \mathbb{E} \left\{ \max_{j \leq k} |\hat{\mathbf{w}}_i^j - \bar{\mathbf{w}}_i^j|^q \right\}^{1/q}. \quad (\text{D.9})$$

Theorem C.5 yields, for $p \in \mathbb{N}$,

$$\widehat{\mathcal{D}}_{2p}(k) \leq C_* L r e^{C_* p L(k\varepsilon)} (d^2 \varepsilon \log(1/\varepsilon))^{1/4}. \quad (\text{D.10})$$

Note that $\mathbf{w}_i^k, \bar{\mathbf{w}}_i^k$ take the form

$$\mathbf{w}_i^k = \mathbf{w}_i^0 + \mathbf{M}_i^k + \mathbf{V}_i^k + \boldsymbol{\varphi}_i^k, \quad (\text{D.11})$$

$$\bar{\mathbf{w}}_i^k = \mathbf{w}_i^0 + \bar{\mathbf{M}}_i^k + \bar{\mathbf{V}}_i^k + \bar{\boldsymbol{\varphi}}_i^k, \quad (\text{D.12})$$

where $\mathbf{M}_i^k, \bar{\mathbf{M}}_i^k$ are martingales with respect to the filtration $\mathcal{F}_k$: $\mathbb{E}\{\mathbf{M}_i^k | \mathcal{F}_{k-1}\} = \mathbf{M}_i^{k-1}$, $\mathbb{E}\{\bar{\mathbf{M}}_i^k | \mathcal{F}_{k-1}\} = \bar{\mathbf{M}}_i^{k-1}$, and $\mathbf{V}_i^k, \bar{\mathbf{V}}_i^k$ are $\mathcal{F}_{k-1}$ -measurable. Explicitly

$$\mathbf{M}_i^k = \sum_{\ell=0}^{k-1} (\mathbf{F}_i(\mathbf{z}_{\ell+1}; \mathbf{w}^\ell) - \mathbb{E}\{\mathbf{F}_i(\mathbf{z}_{\ell+1}; \mathbf{w}^\ell) | \mathcal{F}_\ell\}), \quad (\text{D.13})$$

$$\bar{\mathbf{M}}_i^k = \sum_{\ell=0}^{k-1} \sqrt{2\tau\varepsilon} \mathbf{g}_i^{\ell+1}, \quad (\text{D.14})$$

$$\mathbf{V}_i^k = \sum_{\ell=0}^{k-1} \mathbb{E}\{\mathbf{F}_i(\mathbf{z}_{\ell+1}; \mathbf{w}^\ell) | \mathcal{F}_\ell\} = \sum_{\ell=0}^{k-1} \varepsilon \mathbf{G}(\mathbf{w}_i^\ell; \rho_\ell^{(N)}), \quad (\text{D.15})$$

$$\bar{\mathbf{V}}_i^k = \sum_{\ell=0}^{k-1} \varepsilon \mathbf{G}(\bar{\mathbf{w}}_i^\ell; \rho_{\ell\varepsilon}). \quad (\text{D.16})$$

Finally, $\boldsymbol{\varphi}_i^k, \bar{\boldsymbol{\varphi}}_i^k$ are corrections to satisfy the constraint $\mathbf{w}_i^k, \bar{\mathbf{w}}_i^k \in D$. Indeed the above can be viewed as Skorokhod problems with unknowns $(\mathbf{w}_i, \boldsymbol{\varphi}_i)$ and $(\bar{\mathbf{w}}_i, \bar{\boldsymbol{\varphi}}_i)$.

Using [Slo94, Theorem 1] (where we can set $C_p = (C_* p)^{2p}$ which is the tight constant in the Burkholder-Davis-Gundy inequality), we get

$$\mathcal{D}_{2p}(k) \leq C_* p \left\{ \mathbb{E}([\mathbf{M}_i - \bar{\mathbf{M}}_i]_k^p)^{1/(2p)} + \mathbb{E}(|\mathbf{V}_i - \bar{\mathbf{V}}_i|_k^{2p})^{1/(2p)} \right\}, \quad (\text{D.17})$$

where $[\mathbf{M}]_k$ denotes the quadratic variation of the martingale $\mathbf{M}$, and $|\mathbf{V}|_k$ is the total variation

of the process $\mathbf{V}$. We then have

$$A_{1,p}(k) \equiv \mathbb{E}([\mathbf{M}_i - \overline{\mathbf{M}}_i]_k^p)^{1/(2p)} = \mathbb{E} \left\{ \left[\sum_{\ell=0}^{k-1} |\mathbf{Z}_i^\ell|^2 \right]^p \right\}^{1/(2p)}$$

$$\mathbf{Z}_i^\ell = \varepsilon \nabla \sigma(\mathbf{x}_{\ell+1}; \mathbf{w}_i^k) \left(y_{\ell+1} - \frac{1}{N} \sum_{i=1}^N \sigma(\mathbf{x}_{\ell+1}; \mathbf{w}_i^\ell) \right) + \varepsilon \mathbf{G}(\mathbf{w}_i^\ell; \rho_\ell^{(N)}), \quad (\text{D.18})$$

Note that under the stated assumption the martingale increments $\mathbf{Z}_i^\ell$ are sub-Gaussian with variance proxy upper bounded by $v^2 = C_* \varepsilon^2 \sigma_\infty^2 \gamma^2$. Therefore, by using the moment generating function of χ_d^2 distribution, we have

$$\mathbb{E} \exp \left\{ \frac{\alpha^2}{2dv^2} |\mathbf{Z}_i^\ell|^2 \right\} \leq \left(1 - \frac{\alpha^2}{d} \right)^{-d/2}. \quad (\text{D.19})$$

Hence,

$$\mathbb{E} \exp \left\{ \frac{\alpha^2}{2dv^2} \sum_{\ell=0}^{k-1} |\mathbf{Z}_i^\ell|^2 \right\} \leq \left(1 - \frac{\alpha^2}{d} \right)^{-dk/2}. \quad (\text{D.20})$$

By using the inequality $x^p \leq e^x p!$, this implies, for $\alpha \leq \sqrt{d/2}$,

$$\mathbb{E} \left\{ \left[\frac{\alpha^2}{2dv^2} \sum_{\ell=0}^{k-1} |\mathbf{Z}_i^\ell|^2 \right]^p \right\} \leq p! \left(1 - \frac{\alpha^2}{d} \right)^{-dk/2} \leq p^p e^{\alpha^2 k}. \quad (\text{D.21})$$

Equivalently,

$$\mathbb{E} \left\{ \left[\sum_{\ell=0}^{k-1} |\mathbf{Z}_i^\ell|^2 \right]^p \right\}^{1/(2p)} \leq \left(\frac{\sqrt{2dv}}{\alpha} \right) \sqrt{p} e^{\alpha^2 k/(2p)}. \quad (\text{D.22})$$

By taking $\alpha = \sqrt{p/k}$ (which is allowed provided $p \leq \sqrt{kd/2}$), we obtain that

$$A_{1,p}(k) \leq 10\sqrt{kd}v \leq C_* \sqrt{kd} \varepsilon \sigma_\infty \gamma. \quad (\text{D.23})$$

We next consider the total variation of the process $\mathbf{V}_i$ in Eq. (D.17). We have

$$\begin{aligned} \mathbb{E}(|\mathbf{V}_i - \overline{\mathbf{V}}_i|_k^{2p})^{1/(2p)} &= \mathbb{E} \left\{ \left(\sum_{\ell=0}^{k-1} \varepsilon \left| \mathbf{G}(\mathbf{w}_i^\ell; \rho_\ell^{(N)}) - \mathbf{G}(\overline{\mathbf{w}}_i^\ell; \rho_{\ell\varepsilon}) \right| \right)^{2p} \right\}^{1/(2p)} \\ &\leq \mathbb{E} \left\{ \left(\sum_{\ell=0}^{k-1} \varepsilon \left| \mathbf{G}(\mathbf{w}_i^\ell; \rho_\ell^{(N)}) - \mathbf{G}(\overline{\mathbf{w}}_i^\ell; \rho_\ell^{(N)}) \right| \right)^{2p} \right\}^{1/(2p)} \\ &\quad + \mathbb{E} \left\{ \left(\sum_{\ell=0}^{k-1} \varepsilon \left| \mathbf{G}(\overline{\mathbf{w}}_i^\ell; \rho_\ell^{(N)}) - \mathbf{G}(\overline{\mathbf{w}}_i^\ell; \rho_{\ell\varepsilon}) \right| \right)^{2p} \right\}^{1/(2p)} \\ &\equiv A_{2,p}(k) + A_{3,p}(k). \end{aligned}$$

Using the Lipschitz property of ∇V , ∇U , we get

$$\begin{aligned} A_{2,p}(k) &\leq L\varepsilon \mathbb{E} \left\{ \left(\sum_{\ell=0}^{k-1} |\mathbf{w}_i^\ell - \bar{\mathbf{w}}_i^\ell| \right)^{2p} \right\}^{1/(2p)} \\ &\leq L\varepsilon \sum_{\ell=0}^{k-1} \mathcal{D}_{2p}(\ell). \end{aligned} \quad (\text{D.24})$$

For the second term, we get, by triangular inequality,

$$A_{3,p}(k) \leq \varepsilon \sum_{\ell=0}^{k-1} \mathbb{E} \left\{ \left| \mathbf{G}(\bar{\mathbf{w}}_i^\ell; \rho_\ell^{(N)}) - \mathbf{G}(\bar{\mathbf{w}}_i^\ell; \rho_{\ell\varepsilon}) \right|^{2p} \right\}^{1/(2p)}. \quad (\text{D.25})$$

We next use the expression $\mathbf{G}(\mathbf{w}; \rho) = \nabla V(\mathbf{w}) + \int \nabla U(\mathbf{w}, \mathbf{w}') \rho(d\mathbf{w}')$, and the fact that $\hat{\mathbf{w}}_i^\ell \sim \rho_{\ell\varepsilon}$, to get

$$\begin{aligned} A_{3,p}(k) &\leq \varepsilon \sum_{\ell=0}^{k-1} \mathbb{E} \left\{ \left| \frac{1}{N} \sum_{j=1}^N [\nabla U(\bar{\mathbf{w}}_i^\ell, \mathbf{w}_j^\ell) - \mathbb{E}_{\hat{\mathbf{w}}_j^\ell} \nabla U(\bar{\mathbf{w}}_i^\ell, \hat{\mathbf{w}}_j^\ell)] \right|^{2p} \right\}^{1/(2p)} \\ &\leq \varepsilon \sum_{\ell=0}^{k-1} \mathbb{E} \left\{ \left| \frac{1}{N} \sum_{j=1}^N [\nabla U(\bar{\mathbf{w}}_i^\ell, \mathbf{w}_j^\ell) - \nabla U(\bar{\mathbf{w}}_i^\ell, \hat{\mathbf{w}}_j^\ell)] \right|^{2p} \right\}^{1/(2p)} \\ &\quad + \varepsilon \sum_{\ell=0}^{k-1} \mathbb{E} \left\{ \left| \frac{1}{N} \sum_{j=1}^N [\nabla U(\bar{\mathbf{w}}_i^\ell, \hat{\mathbf{w}}_j^\ell) - \mathbb{E}_{\hat{\mathbf{w}}_j^\ell} \nabla U(\bar{\mathbf{w}}_i^\ell, \hat{\mathbf{w}}_j^\ell)] \right|^{2p} \right\}^{1/(2p)} \\ &\equiv A_{3,p}^{(1)}(k) + A_{3,p}^{(2)}(k). \end{aligned}$$

Using once more the Lipschitz property of ∇U , and the symmetry of the distributions of $(\mathbf{w}^\ell)_{i \leq N}$, $(\hat{\mathbf{w}}^\ell)_{i \leq N}$ under permutations, we obtain

$$A_{3,p}^{(1)}(k) \leq L\varepsilon \sum_{\ell=0}^{k-1} \mathbb{E} \left\{ \left| \mathbf{w}_j^\ell - \hat{\mathbf{w}}_j^\ell \right|^{2p} \right\}^{1/(2p)} \quad (\text{D.26})$$

$$\leq L\varepsilon \sum_{\ell=0}^{k-1} \mathcal{D}_{2p}(\ell) + L\varepsilon \sum_{\ell=0}^{k-1} \widehat{\mathcal{D}}_{2p}(\ell). \quad (\text{D.27})$$

Finally, $|\nabla U(\bar{\mathbf{w}}_i^\ell, \hat{\mathbf{w}}_j^\ell)| \leq L$ and therefore the vector

$$\mathbf{W} = \frac{1}{N} \sum_{j=1}^N [\nabla U(\bar{\mathbf{w}}_i^\ell, \hat{\mathbf{w}}_j^\ell) - \mathbb{E}_{\hat{\mathbf{w}}_j^\ell} \nabla U(\bar{\mathbf{w}}_i^\ell, \hat{\mathbf{w}}_j^\ell)]$$

is sub-Gaussian, with variance proxy upper bounded by $v^2 = L^2/N$. This implies that $\mathbb{E}\{|\mathbf{W}|^{2p}\}^{1/(2p)} \leq C_* \sqrt{dp} v$, and therefore

$$A_{3,p}^{(2)}(k) \leq C_*(k\varepsilon) L \sqrt{\frac{dp}{N}}. \quad (\text{D.28})$$

Substituting (D.23), (D.24), (D.27), (D.28) in Eq. (D.17), we obtain

$$\begin{aligned}\mathcal{D}_{2p}(k) &\leq C_* L p \varepsilon \sum_{\ell=0}^{k-1} \mathcal{D}_{2p}(\ell) + C_* L p(k\varepsilon) \widehat{\mathcal{D}}_{2p}(k) \\ &\quad + C_* p \sqrt{k d} \varepsilon \sigma_\infty \gamma + C_* L p(k\varepsilon) \sqrt{\frac{dp}{N}}.\end{aligned}\tag{D.29}$$

Using Eq. (D.10) and Gronwall inequality, along with the fact that $k\varepsilon \leq T$, this yields

$$\mathcal{D}_{2p}(T/\varepsilon) \leq C_* e^{C_* p L T} \left[\sqrt{\frac{d}{N}} \vee (\sigma_\infty \gamma \sqrt{d\varepsilon}) \vee \left(L r (d^2 \varepsilon \log(1/\varepsilon))^{1/4} \right) \right].$$

By using Eq. (D.10) again, we get

$$\begin{aligned}\mathcal{D}_{2p}(T/\varepsilon) + \widehat{\mathcal{D}}_{2p}(T/\varepsilon) &\leq C_* e^{C_* p L T} \left[\sqrt{\frac{d}{N}} \vee (\sigma_\infty \gamma \sqrt{d\varepsilon}) \vee \left(L r (d^2 \varepsilon \log(1/\varepsilon))^{1/4} \right) \right] \\ &\equiv e^{C_* p L T} \text{err}(N, d, \varepsilon).\end{aligned}\tag{D.30}$$

By Markov inequality along with the Jensen inequality applied to the convex function x^{2p} , we have

$$\begin{aligned}\mathbb{P} \left\{ \frac{1}{N} \sum_{i=1}^N |\mathbf{w}_i^k - \hat{\mathbf{w}}_i^k| \geq \Delta \right\} &\leq \frac{1}{\Delta^{2p}} \mathbb{E} \left\{ \left[\frac{1}{N} \sum_{i=1}^N |\mathbf{w}_i^k - \hat{\mathbf{w}}_i^k| \right]^{2p} \right\} \\ &\leq \frac{1}{\Delta^{2p}} \mathbb{E} \left\{ \frac{1}{N} \sum_{i=1}^N |\mathbf{w}_i^k - \hat{\mathbf{w}}_i^k|^{2p} \right\} \\ &\leq \frac{1}{\Delta^{2p}} \left[\mathcal{D}_{2p}(T/\varepsilon) + \widehat{\mathcal{D}}_{2p}(T/\varepsilon) \right]^{2p} \\ &\leq \frac{1}{\Delta^{2p}} e^{2C_* p^2 L T} \text{err}(N, d, \varepsilon)^{2p},\end{aligned}$$

where in the third step we used (D.8) and (D.9). Set $\Delta = z e^{C_* p L T} \text{err}(N, d, \varepsilon)$. Thus, we obtain

$$\frac{1}{N} \sum_{i=1}^N |\mathbf{w}_i^k - \hat{\mathbf{w}}_i^k| \leq z e^{C_* p L t} \text{err}(N, d, \varepsilon),\tag{D.31}$$

with probability at least $1 - z^{-2p}$.

The bounds in Eq. (D.3) follow straightforwardly from Eq. (D.31) as in the proofs of Lemma 3.3 and 3.4 in the supplementary material of [MMN18]. $\square$

E Regularity of the solutions of the PDE (3.9) ($\delta > 0$)

In this section we prove some standard regularity properties of the solutions of the PDE (3.9), for $\delta > 0$, and indeed for the more general PDE (B.2). First of all, we show that the weak solution of

the PDE (B.2) is in fact strong, i.e., $\rho \in \mathcal{C}^{2,1}(\Omega^\delta, [0, T])$ and the equation (B.2) holds pointwise. We will then prove upper bounds on $\nabla K^\delta * \rho$ and $\nabla U^\delta * \rho$ that are uniform in δ . These will be crucial in order to take the $\delta \rightarrow 0$ limit in the next section.

We start by proving a bound on the $\mathcal{L}^\infty$ norm of ρ . In the proofs of the two lemmas that follow, we assume without loss of generality that $\tau = 1$.

Lemma E.1 (Bound on $\mathcal{L}^\infty$ norm). *Let ρ_t be a weak solution of the PDE (B.2) with initial and boundary conditions (B.3). Recall that ρ_t has a density with respect to Lebesgue measure, denoted by $\rho(\cdot, t)$. Then, there exists a constant $C(\Omega)$ such that, by letting $L = (\|\nabla V\|_{\mathcal{L}^\infty(\Omega)} \vee \|\nabla U\|_{\mathcal{L}^\infty(\Omega \times \Omega)})$, we have*

$$\|\rho(\cdot, t)\|_{\mathcal{L}^\infty(\Omega)} \leq \|\rho_{\text{init}}\|_{\mathcal{L}^\infty(\Omega)} e^{C(\Omega)L^2 t}. \quad (\text{E.1})$$

Proof. Any solution the PDE (B.2) satisfies Eq. (B.6). Given a measurable (Borel) function $\rho \in m\mathcal{B}(\Omega \times [0, T])$, denote by $\mathcal{D}(\rho) \in m\mathcal{B}(\Omega \times [0, T])$ the function given by the right-hand side of (B.2). Let $C(\Omega)$ be the constant in the statement of Theorem G.1 (part 3) and let $C_{U,V} \equiv C(\Omega)(\|\nabla V\|_{\mathcal{L}^\infty(\Omega)} + \|\nabla U\|_{\mathcal{L}^\infty(\Omega \times \Omega)})$. We then have

$$\begin{aligned} \|\mathcal{D}(\rho)(\cdot, t)\|_{\mathcal{L}^\infty(\Omega)} &\leq \|\rho_{\text{init}}\|_{\mathcal{L}^\infty(\Omega)} + (\|\nabla V\|_{\mathcal{L}^\infty(\Omega)} + \|\nabla U\|_{\mathcal{L}^\infty(\Omega \times \Omega)}) \\ &\quad \int_0^t \sup_{\mathbf{x} \in \Omega} \|\nabla G^\Omega(\mathbf{x}, \cdot; t-s)\|_{\mathcal{L}^1(\Omega)} \|\rho(\cdot, s)\|_{\mathcal{L}^\infty(\Omega)} ds \\ &\leq \|\rho_{\text{init}}\|_{\mathcal{L}^\infty(\Omega)} + C_{U,V} \int_0^t (t-s)^{-1/2} \|\rho(\cdot, s)\|_{\mathcal{L}^\infty(\Omega)} ds. \end{aligned} \quad (\text{E.2})$$

Hence

$$\|\mathcal{D}(\rho)\|_{\mathcal{L}^\infty(\Omega \times [0, T])} \leq \|\rho_{\text{init}}\|_{\mathcal{L}^\infty(\Omega)} + C_{U,V} \sqrt{T} \|\rho\|_{\mathcal{L}^\infty(\Omega \times [0, T])}. \quad (\text{E.3})$$

Proceeding analogously for two different densities $\rho, \tilde{\rho}$, we get

$$\|\mathcal{D}(\rho) - \mathcal{D}(\tilde{\rho})\|_{\mathcal{L}^\infty(\Omega \times [0, T])} \leq C_{U,V} \sqrt{T} \|\rho - \tilde{\rho}\|_{\mathcal{L}^\infty(\Omega \times [0, T])}. \quad (\text{E.4})$$

Hence $\mathcal{D}$ maps $\mathcal{L}^\infty(\Omega \times [0, T])$ into itself, and is a contraction for $C_{U,V} \sqrt{T} < 1$. Therefore, it must have a unique fixed point in $\mathcal{L}^\infty$ that coincides with the unique solution of PDE (B.2). Let $T_0 = 1/(4C_{U,V}^2)$. Then for that fixed point $\rho \in \mathcal{L}^\infty(\Omega \times [0, T])$ we have from Eq. (E.3)

$$\|\rho\|_{\mathcal{L}^\infty(\Omega \times [0, T_0])} \leq 2 \|\rho_{\text{init}}\|_{\mathcal{L}^\infty(\Omega)}. \quad (\text{E.5})$$

The desired claim follow by iterating this inequality $\lceil t/T_0 \rceil$ times. $\square$

Lemma E.2 (Strong solutions of PDE). *Let ρ_t be a weak solution of the PDE (B.2) with initial and boundary conditions (B.3), and recall that, for any $t \leq T < \infty$, this has a density $\rho(\cdot, t)$, with $\rho \in \mathcal{L}^\infty(\Omega \times [0, T])$. Fix $q \in \mathbb{N}$. If $\rho_{\text{init}} \in \mathcal{C}^q(\Omega)$, then $\rho \in \mathcal{C}^{q,1}(\Omega, [0, T])$.*

Proof. We prove the claim for $q = 2$. For larger values of q , the proof is similar and it only requires to iterate the argument.

The proof uses the same bootstrap technique of [MMN18][Supplementary material, Lemma 6.7]. The only difference is that the Duhamel formula of Eq. (B.6) involves the Neumann heat kernel in Ω instead of the heat kernel in $\mathbb{R}^d$.

Let $S = \Omega \times [0, T]$ and, for $u : \Omega \times [0, T] \rightarrow \mathbb{R}$. For $r \in \mathbb{N}$, $\alpha = (\alpha_1, \dots, \alpha_d) \in \mathbb{N}^d$, let $D_t^r D_{\mathbf{x}}^\alpha u$ be the generalized derivative of u , and define the parabolic seminorm

$$\langle\langle u \rangle\rangle_{\mathcal{L}^p(S)}^{(j)} \equiv \sum_{|\alpha|+2r=j} \|D_t^r D_{\mathbf{x}}^\alpha u\|_{\mathcal{L}^p(S)}. \quad (\text{E.6})$$

The proof of [MMN18][Supplementary material, Lemma 6.7] uses the following inequality from [LSU88][Chapter IV, Section 3, Eq. (3.1)]

$$\langle\langle G *_2 u \rangle\rangle_{\mathcal{L}^p(S)}^{(2m+2)} \leq C \langle\langle u \rangle\rangle_{\mathcal{L}^p(S)}^{(2m)}, \quad (\text{E.7})$$

$$G *_2 u(\mathbf{x}, t) \equiv \int_{\mathbb{R}^d} \int_0^t G(\mathbf{x}, \mathbf{y}, t-s) u(\mathbf{y}, s) d\mathbf{y} ds. \quad (\text{E.8})$$

Furthermore, (G.11) of Theorem G.1 yields

$$\langle\langle G^\Omega *_2 u \rangle\rangle_{\mathcal{L}^p(S)}^{(2m+2)} \leq \langle\langle G *_2 u \rangle\rangle_{\mathcal{L}^p(S)}^{(2m+2)} + \langle\langle G_R^\Omega *_2 u \rangle\rangle_{\mathcal{L}^p(S)}^{(2m+2)}. \quad (\text{E.9})$$

Since $G_R^\Omega \in \mathcal{C}^\infty(\Omega \times \Omega \times [0, T])$, we have that

$$\langle\langle G_R^\Omega *_2 u \rangle\rangle_{\mathcal{L}^p(S)}^{(2m+2)} \leq C \|u\|_{\mathcal{L}^p(S)},$$

which immediately implies that

$$\langle\langle G^\Omega *_2 u \rangle\rangle_{\mathcal{L}^p(S)}^{(2m+2)} \leq C \langle\langle u \rangle\rangle_{\mathcal{L}^p(S)}^{(2m)} + C \|u\|_{\mathcal{L}^p(S)}. \quad (\text{E.10})$$

The proof of [MMN18][Supplementary material, Lemma 6.7] can be repeated verbatimly with (E.7) replaced by (E.10). $\square$

As a consequence of the last lemma, the PDE (B.2) admits unique strong solutions $\rho \in \mathcal{C}^{2,1}(\Omega, [0, T])$ with initial condition ρ_{init} and Neumann boundary condition. We will use $\rho(t)$ as shortcut for $\rho(\cdot, t)$. The rest of this appendix is devoted to prove further regularity results for $\rho(t)$, which will be crucial in the proofs provided in Appendix F. To emphasize the dependence of ρ on δ , we will denote this solution by ρ^δ .

In what follows, we will set the initial condition $\rho^\delta(0) \equiv \rho_{\text{init}}^\delta$ at $\delta > 0$ to be defined via $\rho_{\text{init}}^\delta(\mathbf{w}) = \lambda_\delta^{-d} \rho_{\text{init}}(\mathbf{w}/\lambda_\delta)$, with λ_δ given by Eq. (3.4)

It is useful to recall the definition of free energy, which is given by

$$\begin{aligned} F^\delta(\rho) &= \frac{1}{2} R^\delta(\rho) - \tau S(\rho) \\ &= \frac{\nu_0}{2} \|f - K^\delta * \rho\|_{\mathcal{L}^2(\Omega)}^2 + \tau \int \rho(\mathbf{x}) \log \rho(\mathbf{x}) d\mathbf{x}. \end{aligned} \quad (\text{E.11})$$

The following lemma provides an expression for the derivative of the free energy with respect to time. Such an expression immediately yields an upper bound on the $\mathcal{L}^2(\Omega)$ norm of $K^\delta * \rho^\delta(t)$ which is independent of δ .

Lemma E.3. *Let $\rho^\delta \in \mathcal{C}^{2,1}(\Omega^\delta, [0, T])$ be the solution of the PDE (B.2) with initial and boundary conditions (B.3). Then,*

$$\frac{d}{dt} F^\delta(\rho^\delta(t)) = - \int |\nabla(\Psi(\mathbf{x}; \rho^\delta(t)) + \tau \log \rho^\delta(\mathbf{x}, t))|^2 \rho^\delta(\mathbf{x}, t) d\mathbf{x}. \quad (\text{E.12})$$

Proof. By definition

$$\begin{aligned} F^\delta(\rho^\delta(t)) &= \frac{\nu_0}{2} \|f\|_{\mathcal{L}^2(\Omega)}^2 + \int V(\mathbf{w}) \rho^\delta(\mathbf{w}, t) d\mathbf{w} \\ &\quad + \frac{1}{2} \int U(\mathbf{w}_1 - \mathbf{w}_2) \rho^\delta(\mathbf{w}_1, t) d\mathbf{w}_1 \rho^\delta(\mathbf{w}_2, t) d\mathbf{w}_2 \\ &\quad + \tau \int \rho^\delta(\mathbf{w}, t) \log \rho^\delta(\mathbf{w}, t) d\mathbf{w}. \end{aligned}$$

By differentiating $F^\delta(\rho^\delta(t))$ along the solution of (B.2), we obtain

$$\begin{aligned} \frac{d}{dt} F^\delta(\rho^\delta(t)) &= \int V(\mathbf{w}) \partial_t \rho^\delta(\mathbf{w}, t) d\mathbf{w} \\ &\quad + \int U(\mathbf{w}_1 - \mathbf{w}_2) \rho^\delta(\mathbf{w}_1, t) \partial_t \rho^\delta(\mathbf{w}_2, t) d\mathbf{w}_1 d\mathbf{w}_2 \\ &\quad + \tau \int (1 + \log \rho^\delta(\mathbf{w}, t)) \partial_t \rho^\delta(\mathbf{w}, t) d\mathbf{w} \\ &= \int (\Psi(\mathbf{w}, \rho^\delta(t)) + \tau \log \rho^\delta(\mathbf{w}, t) + \tau) \partial_t \rho^\delta(\mathbf{w}, t) d\mathbf{w} \\ &= \int (\Psi(\mathbf{w}, \rho^\delta(t)) + \tau \log \rho^\delta(\mathbf{w}, t) + \tau) \\ &\quad \nabla \cdot (\rho^\delta(\mathbf{w}, t) \nabla (\Psi(\mathbf{w}, \rho^\delta(t)) + \tau \log \rho^\delta(\mathbf{w}, t))) d\mathbf{w} \\ &= - \int \langle \nabla (\Psi(\mathbf{w}, \rho^\delta(t)) + \tau \log \rho^\delta(\mathbf{w}, t)), \\ &\quad \nabla (\Psi(\mathbf{w}, \rho^\delta(t)) + \tau \log \rho^\delta(\mathbf{w}, t)) \rangle \rho^\delta(\mathbf{w}, t) d\mathbf{w} \\ &= - \int |\nabla (\Psi(\mathbf{w}; \rho^\delta(t)) + \tau \log \rho^\delta(\mathbf{w}, t))|^2 \rho^\delta(\mathbf{w}, t) d\mathbf{w} \end{aligned} \tag{E.13}$$

□

Corollary E.4. Let $\rho^\delta \in \mathcal{C}^{2,1}(\Omega^\delta, [0, T])$ be the solution of the PDE (B.2) with initial and boundary conditions (B.3). Then,

$$\nu_0 \|K^\delta * \rho^\delta(t) - f\|_{\mathcal{L}^2(\Omega)}^2 \leq 2F^\delta(\rho^\delta(0)) + 2\tau \log |\Omega^\delta|, \tag{E.14}$$

where $|\Omega^\delta|$ denotes the volume of the set Ω^δ .

Proof. By Lemma E.3 we have $F^\delta(\rho^\delta(t)) \leq F^\delta(\rho^\delta(0))$. The claim follows by substituting the definition of $F^\delta(\rho^\delta)$ and using $S(\rho^\delta) \leq \log |\Omega^\delta|$. □

Remark E.1. By Corollary E.4, we are able to provide a δ -free upper bound on $\nu_0 \|K^\delta * \rho^\delta(t) - f\|_{\mathcal{L}^2(\Omega)}^2$. Specifically, $\Omega^\delta \subseteq \Omega$ and hence $|\Omega^\delta| \leq |\Omega|$. We also have

$$F^\delta(\rho^\delta(0)) \leq \nu_0 \|f\|_{\mathcal{L}^2(\Omega)}^2 + \nu_0 \|K^\delta * \rho^\delta(0)\|_{\mathcal{L}^2(\Omega)} - \tau S(\rho^\delta(0)).$$

Note that

$$S(\rho^\delta(0)) = S(\rho_{\text{init}}) + d \log \lambda_\delta.$$

Since $\lambda_\delta \rightarrow 1$ as $\delta \rightarrow 0$, there exists a $C_* > 0$ such that for $\delta < C_*$, $\lambda_\delta \geq 1/2$. Thus, the term $S(\rho^\delta(0))$ has a δ -free upper bound.

By Young's inequality it only remains to give a δ -free upper bound on the quantity $\|\rho^\delta(0)\|_{\mathcal{L}^2(\Omega)}$. Let us write

$$\begin{aligned}\|\rho^\delta(0)\|_{\mathcal{L}^2(\Omega)}^2 &\leq \int_{\Omega^\delta} \lambda_\delta^{-2d} \rho_{\text{init}}^2(\mathbf{w}/\lambda_\delta) d\mathbf{w} \\ &= \int_{\Omega} \lambda_\delta^{-2d} \rho_{\text{init}}^2(\mathbf{x}) \lambda_\delta^d d\mathbf{x} = \lambda_\delta^{-d} \|\rho_{\text{init}}\|_{\mathcal{L}^2(\Omega)}^2.\end{aligned}$$

Again, for $\delta < C_*$, $\lambda_\delta \geq 1/2$. Also, by Assumption (A5) and the fact that Ω is compact, we have $\|\rho_{\text{init}}\|_{\mathcal{L}^2(\Omega)}^2 < \infty$, which concludes the claim.

We next prove δ -free upper bound on the gradient of $\nabla K^\delta * \rho^\delta$.

Lemma E.5. *Let $\rho^\delta \in \mathcal{C}^{2,1}(\Omega^\delta, [0, T])$ be the solution of the PDE (B.2) with initial and boundary conditions (B.3). Then, the following bound holds:*

$$\begin{aligned}\int_0^T \int |\nabla U * \rho^\delta(\mathbf{x}, t)|^2 \rho^\delta(\mathbf{x}, t) d\mathbf{x} dt + 2\tau \int_0^T \int |\nabla(K^\delta * \rho^\delta)(\mathbf{x}, t)|^2 d\mathbf{x} dt \\ \leq T \|\nabla V\|_{\mathcal{L}^\infty(\Omega)}^2 + 2\nu_0 \|f\|_{\mathcal{L}^2(\Omega)}^2 + 4F^\delta(\rho^\delta(0)) + 4\tau \log |\Omega^\delta|.\end{aligned}\tag{E.15}$$

Proof. Denote by $\langle f, g \rangle = \int f(\mathbf{x})g(\mathbf{x})d\mathbf{x}$ the standard scalar product in $\mathcal{L}^2$. Then,

$$\begin{aligned}\frac{1}{2} \frac{d}{dt} \langle U * \rho^\delta(t), \rho^\delta(t) \rangle &= \langle U * \rho^\delta(t), \partial_t \rho^\delta(t) \rangle \\ &= \langle U * \rho^\delta(t), \nabla \cdot (\rho^\delta(t) \nabla (V + U * \rho^\delta(t))) + \tau \Delta \rho^\delta(t) \rangle \\ &= \langle U * \rho^\delta(t), \nabla \cdot (\rho^\delta(t) \nabla V) \rangle + \langle U * \rho^\delta(t), \nabla \cdot (\rho^\delta(t) \nabla (U * \rho^\delta(t))) \rangle \\ &\quad + \tau \langle U * \rho^\delta(t), \Delta \rho^\delta(t) \rangle \\ &= - \int \langle \nabla (U * \rho^\delta)(\mathbf{x}, t), \nabla V(\mathbf{x}) \rangle \rho^\delta(\mathbf{x}, t) d\mathbf{x} \\ &\quad - \int |\nabla (U * \rho^\delta)(\mathbf{x}, t)|^2 \rho^\delta(\mathbf{x}, t) d\mathbf{x} - \tau \int |\nabla(K^\delta * \rho^\delta)(\mathbf{x}, t)|^2 d\mathbf{x} \\ &\leq \frac{1}{2} \int |\nabla V(\mathbf{x})|^2 \rho^\delta(\mathbf{x}, t) d\mathbf{x} - \frac{1}{2} \int |\nabla (U * \rho^\delta)(\mathbf{x}, t)|^2 \rho^\delta(\mathbf{x}, t) d\mathbf{x} \\ &\quad - \tau \int |\nabla(K^\delta * \rho^\delta)(\mathbf{x}, t)|^2 d\mathbf{x}.\end{aligned}\tag{E.16}$$

By integrating (E.16) between 0 and T , we obtain

$$\begin{aligned}\int_0^T \int |\nabla U * \rho^\delta(\mathbf{x}, t)|^2 \rho^\delta(\mathbf{x}, t) d\mathbf{x} dt + 2\tau \int_0^T \int |\nabla(K^\delta * \rho^\delta)(\mathbf{x}, t)|^2 d\mathbf{x} dt \\ \leq T \|\nabla V\|_\infty^2 - \langle \rho^\delta(T), U * \rho^\delta(T) \rangle + \langle \rho^\delta(0), U * \rho^\delta(0) \rangle \\ \leq T \|\nabla V\|_{\mathcal{L}^\infty(\Omega)}^2 + \nu_0 \|K^\delta * \rho^\delta(0)\|_{\mathcal{L}^2(\Omega)}^2.\end{aligned}\tag{E.17}$$

Hence, (E.15) follows from Corollary E.4. $\square$

Remark E.2. Note that by virtue of Lemma E.5, we are able to get a δ -free upper bound on the left-hand side of (E.15). Indeed, by definition of ∇V as per (B.4) and using Assumption (A3), we have the δ -free bound:

$$\|\nabla V\|_{\mathcal{L}^\infty(\Omega)} \leq \nu_0 \|K^\delta\|_{\mathcal{L}^1(\Omega)} \|\nabla f\|_{\mathcal{L}^\infty(\Omega)} = \|\nabla f\|_{\mathcal{L}^\infty(\Omega)} < C_*. \quad (\text{E.18})$$

In addition, by Remark E.1, $\|K^\delta * \rho^\delta(0)\|_{\mathcal{L}^2(\Omega)}^2$ has δ -free bound.

F Global convergence: Proof of Theorems 5.2 and 5.3

We start by showing that ρ^δ admits a limit in a suitable functional space as $\delta \rightarrow 0$.

Lemma F.1 (Existence of converging subsequence). *Let $\rho^\delta \in \mathcal{C}^{2,1}(\Omega^\delta, [0, T])$ be the unique solution of the PDE (B.2) with initial and boundary conditions (B.3). Then, the family $(\rho^\delta)_{\delta>0}$ is relatively compact in the space $\mathcal{C}([0, T], \mathcal{P}_2(\Omega))$. In particular any sequence $(\rho^{\delta_n})_{n \geq 1}$, admits a converging subsequence.*

Proof. This follows from the Ascoli-Arzelà's theorem. Notice that $\mathcal{P}_2(\Omega)$ is compact due to the compactness of Ω . Therefore, it is sufficient to prove that the family is equicontinuous. Using the representation in terms of nonlinear dynamics (cf. Appendix C), we have

$$W_2(\rho_t, \rho_s)^2 \leq \mathbb{E}\{|X_t - X_s|^2\}. \quad (\text{F.1})$$

Note that we omit for simplicity the dependence on δ . Recall that the nonlinear dynamic satisfies (for $\mathbf{b}(\mathbf{x}, t) \equiv -\nabla \Psi(\mathbf{x}, \rho_t)$)

$$X_t = \int_0^t \mathbf{b}(X_r, r) dr + \sqrt{2\tau} B_t + \Phi_t \quad (\text{F.2})$$

$$\equiv V_t + \sqrt{2\tau} B_t + \Phi_t. \quad (\text{F.3})$$

By [Slo01, Theorem 2.2], we have

$$\mathbb{E}\{|X_t - X_s|^2\} \leq C_* \tau \mathbb{E}\{[B]_s^t\} + C_* \mathbb{E}\{(|V|_s^t)^2\}. \quad (\text{F.4})$$

where $[B]_s^t$ denotes the quadratic variation of B , and $|V|_s^t$ the total variation of V between times s and t . We thus have

$$\begin{aligned} W_2(\rho_t, \rho_s)^2 &\leq \mathbb{E}\{|X_t - X_s|^2\} \leq C_* \tau (t - s) + C_* \mathbb{E}\left\{\left(\int_s^t |\mathbf{b}(X_r, r)| dr\right)^2\right\} \\ &\leq C_* \tau (t - s) + C_* (t - s) \int_s^t \mathbb{E}\{|\mathbf{b}(X_r, r)|^2\} dr. \end{aligned} \quad (\text{F.5})$$

Hence, in order to prove uniform continuity, it is sufficient to show that, for $s, t \leq T$, $\int_s^t \mathbb{E}\{|\mathbf{b}(X_r, r)|^2\} dr \leq C$ where C is bounded uniformly in δ . In order to show that this is the case, notice that

$$\int_s^t \mathbb{E}\{|\mathbf{b}(X_r, r)|^2\} dr = \int_s^t \mathbb{E}\{|\nabla V(X_r) + \nabla U * \rho(X_r, r)|^2\} dr \quad (\text{F.6})$$

$$\leq 2(t - s) \|\nabla V\|_{\mathcal{L}^\infty(\Omega^\delta)}^2 + 2 \int_s^t \int |\nabla U * \rho(\mathbf{x}, r)|^2 \rho(\mathbf{x}, r) d\mathbf{x} dr, \quad (\text{F.7})$$

and the claim follows from Lemma E.5. $\square$

We have now proved that the sequence $(\rho^{\delta_n})_{n \geq 1}$ admits a converging subsequence, where $\delta_n \rightarrow 0$ as $n \rightarrow \infty$. Fix such a convergent subsequence and, with an abuse of notation, also denote it by $(\rho^{\delta_n})_{n \geq 1}$. Let $\rho^\infty \in \mathcal{C}([0, T], \mathcal{P}_2(\Omega))$ be its limit.

Recall that ρ^{δ_n} is supported in Ω^{δ_n} . Hence, $K^{\delta_n} * \rho^{\delta_n}$ is supported in Ω and $K^{\delta_n} * \rho^{\delta_n} \in \mathcal{P}_2(\Omega)$. We will now show that $(K^{\delta_n} * \rho^{\delta_n})_{n \geq 1}$ has the same limit as $(\rho^{\delta_n})_{n \geq 1}$ in $\mathcal{C}([0, T], \mathcal{P}_2(\Omega))$.

Lemma F.2. *The sequence $(K^{\delta_n} * \rho^{\delta_n})_{n \geq 1}$ also converges in $\mathcal{C}([0, T], \mathcal{P}_2(\Omega))$ to ρ^∞ .*

Proof. By Lemma F.1, the result is implied by the following claim:

$$\lim_{n \rightarrow \infty} \sup_{0 \leq t \leq T} W_2(K^{\delta_n} * \rho_t^{\delta_n}, \rho_t^{\delta_n}) = 0. \quad (\text{F.8})$$

Note that, for bounded Ω ,

$$\left(\int |\mathbf{x} - \mathbf{y}|^2 \gamma(d\mathbf{x}, d\mathbf{y}) \right)^{1/2} \leq \text{diam}(\Omega)^{1/2} \left(\int |\mathbf{x} - \mathbf{y}| \gamma(d\mathbf{x}, d\mathbf{y}) \right)^{1/2},$$

for any coupling γ of the probability distributions of $\mathbf{x}$ and $\mathbf{y}$. Hence,

$$W_2(\rho_1, \rho_2) \leq \text{diam}(\Omega)^{1/2} W_1(\rho_1, \rho_2). \quad (\text{F.9})$$

As an application,

$$W_2^2(K^{\delta_n} * \rho_t^{\delta_n}, \rho_t^{\delta_n}) \leq \text{diam}(\Omega)^{1/2} W_1(K^{\delta_n} * \rho_t^{\delta_n}, \rho_t^{\delta_n}). \quad (\text{F.10})$$

Thus, it suffices to show that $\sup_{0 \leq t \leq T} W_1(K^{\delta_n} * \rho_t^{\delta_n}, \rho_t^{\delta_n}) \rightarrow 0$ as $n \rightarrow \infty$.

Note that

$$W_1(K^{\delta_n} * \rho_t^{\delta_n}, \rho_t^{\delta_n}) \leq \mathbb{E}\{|(\mathbf{K}_{\delta_n} + \mathbf{X}_{\delta_n}) - \mathbf{X}_{\delta_n}|\} = \mathbb{E}\{|\mathbf{K}_{\delta_n}|\},$$

where the random variables $\mathbf{K}_{\delta_n}$ and $\mathbf{X}_{\delta_n}$ have distributions K^{δ_n} and $\rho_t^{\delta_n}$, respectively. The quantity $\mathbb{E}\{|\mathbf{K}_{\delta_n}|\}$ is $O(\delta)$, since K has bounded absolute first moment, which completes the proof. $\square$

We will now prove a stronger convergence result.

Lemma F.3 (Convergence in $\mathcal{L}^2$). *The measure ρ^∞ has a density, which is the limit in $\mathcal{L}^2(\Omega \times [0, T])$ of the sequence $(K^{\delta_n} * \rho^{\delta_n})_{n \geq 1}$.*

Proof. By Corollary E.4, we have that, for any $n \geq 1$, $K^{\delta_n} * \rho^{\delta_n} \in \mathcal{L}^2(\Omega \times [0, T])$. Let us show that $(K^{\delta_n} * \rho^{\delta_n})_{n \geq 1}$ is a Cauchy sequence in $\mathcal{L}^2(\Omega \times [0, T])$.

As $K^{\delta_n} * \rho_t^{\delta_n} \in \mathcal{L}^2(\Omega)$ for every $t \in [0, T]$, its Fourier transform exists and we denote it by $\widehat{K^{\delta_n} * \rho_t^{\delta_n}}$. Hence, by applying Parseval's theorem, we have

$$\begin{aligned} & \limsup_{n, n' \rightarrow \infty} \|K^{\delta_n} * \rho^{\delta_n} - K^{\delta_{n'}} * \rho^{\delta_{n'}}\|_{\mathcal{L}^2(\Omega \times [0, T])}^2 \\ &= \limsup_{n, n' \rightarrow \infty} \int_0^T \|K^{\delta_n} * \rho_t^{\delta_n} - K^{\delta_{n'}} * \rho_t^{\delta_{n'}}\|_{\mathcal{L}^2(\Omega)}^2 dt \\ &= \limsup_{n, n' \rightarrow \infty} \int_0^T \int_{\mathbb{R}^d} |\widehat{K^{\delta_n} * \rho_t^{\delta_n}}(\boldsymbol{\lambda}) - \widehat{K^{\delta_{n'}} * \rho_t^{\delta_{n'}}}(\boldsymbol{\lambda})|^2 d\boldsymbol{\lambda} dt. \end{aligned} \quad (\text{F.11})$$

Fix $\Lambda > 1$ and decompose the integral in the right-hand side of (F.11) as

$$\begin{aligned} & \limsup_{n, n' \rightarrow \infty} \int_0^T \int_{|\lambda| < \Lambda} |\widehat{K^{\delta_n} * \rho_t^{\delta_n}(\lambda)} - \widehat{K^{\delta_{n'}} * \rho_t^{\delta_{n'}}(\lambda)}|^2 d\lambda dt \\ & + \limsup_{n, n' \rightarrow \infty} \int_0^T \int_{|\lambda| \geq \Lambda} |\widehat{K^{\delta_n} * \rho_t^{\delta_n}(\lambda)} - \widehat{K^{\delta_{n'}} * \rho_t^{\delta_{n'}}(\lambda)}|^2 d\lambda dt. \end{aligned} \quad (\text{F.12})$$

Consider the first term of (F.12). By Lemma F.2, and since by Jensen's inequality $W_1(\rho_1, \rho_2) \leq W_2(\rho_1, \rho_2)$ for any two distributions ρ_1, ρ_2 , we have $W_1(K^{\delta_n} * \rho_t^{\delta_n} - K^{\delta_{n'}} * \rho_t^{\delta_{n'}}) \rightarrow 0$, as $n, n' \rightarrow \infty$. Since for the complex exponential functions $\|e^{i\langle \lambda, \mathbf{x} \rangle}\|_{\text{Lip}} \leq |\lambda|$, by definition of 1-Wasserstein distance, the integrand in the first term converges pointwise to 0. Furthermore, the integrand is upper bounded by an integrable function, since $|\widehat{K^{\delta_n} * \rho_t^{\delta_n}(\lambda)}| \leq \|K^{\delta_n} * \rho_t^{\delta_n}\|_{\mathcal{L}^2(\Omega)} \leq C$ for all n and every $t \in [0, T]$. Hence, by dominated convergence, the first integral in (F.12) converges to 0.

As for the second term of (F.12), the following chain of inequalities holds:

$$\begin{aligned} & \limsup_{n, n' \rightarrow \infty} \int_0^T \int_{|\lambda| \geq \Lambda} |\widehat{K^{\delta_n} * \rho_t^{\delta_n}(\lambda)} - \widehat{K^{\delta_{n'}} * \rho_t^{\delta_{n'}}(\lambda)}|^2 d\lambda dt \\ & \leq 4 \sup_{n \geq 1} \int_0^T \int_{|\lambda| \geq \Lambda} |\widehat{K^{\delta_n} * \rho_t^{\delta_n}(\lambda)}|^2 d\lambda dt \\ & \leq \frac{4}{\Lambda^2} \sup_{n \geq 1} \int_0^T \int_{|\lambda| \geq \Lambda} |\lambda|^2 |\widehat{K^{\delta_n} * \rho_t^{\delta_n}(\lambda)}|^2 d\lambda dt \\ & \leq \frac{4}{\Lambda^2} \sup_{n \geq 1} \int_0^T \int_{\mathbb{R}^d} |\lambda|^2 |\widehat{K^{\delta_n} * \rho_t^{\delta_n}(\lambda)}|^2 d\lambda dt \\ & = \frac{4}{\Lambda^2} \sup_{n \geq 1} \int_0^T \int_{\mathbb{R}^d} |\nabla \widehat{K^{\delta_n} * \rho_t^{\delta_n}(\lambda)}|^2 d\lambda dt \\ & = \frac{4}{\Lambda^2} \sup_{n \geq 1} \int_0^T \int_{\Omega} |\nabla K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x})|^2 d\mathbf{x} dt, \end{aligned} \quad (\text{F.13})$$

where in the last equality we have applied again Parseval's theorem. By Lemma E.5, the integral in the right-hand side of (F.13) is upper bounded by a constant independent of n . Therefore, as $\Lambda \rightarrow \infty$, the second term of (F.12) converges to 0.

As a result, $(K^{\delta_n} * \rho^{\delta_n})_{n \geq 1}$ is a Cauchy sequence in $\mathcal{L}^2(\Omega \times [0, T])$. Let $\tilde{\rho}^\infty \in \mathcal{L}^2(\Omega \times [0, T])$ be its limit. Furthermore, by Lemma F.2, $(K^{\delta_n} * \rho^{\delta_n})_{n \geq 1}$ has limit ρ^∞ in $\mathcal{C}([0, T], \mathcal{P}_2(\Omega))$. Therefore, the measures $\rho_t^\infty(d\mathbf{x})dt$ and $\tilde{\rho}_t^\infty(\mathbf{x})d\mathbf{x}dt$ coincide. This implies that the measure ρ_t^∞ has for almost every $t \in [0, T]$ the density $\tilde{\rho}_t^\infty \in \mathcal{L}^2(\Omega \times [0, T])$, and the proof is complete. $\square$

From now on, with an abuse of notation, we will use ρ^∞ to denote also the density which is the limit in $\mathcal{L}^2(\Omega \times [0, T])$ of the sequence $(K^{\delta_n} * \rho^{\delta_n})_{n \geq 1}$.

Lemma F.4 (Convergence to a weak solution of the limit PDE). *Let ρ^∞ be the limit in $\mathcal{C}([0, T], \mathcal{P}_2(\Omega))$ of the converging sequence $(\rho^{\delta_n})_{n \geq 1}$. Then, ρ^∞ is a weak solution of the PDE (A.1) with initial and boundary conditions (A.2).*

Proof. By Lemma F.3, we have that $\rho^\infty \in \mathcal{L}^2(\Omega \times [0, T])$. Choose a test function $h \in \mathcal{C}^{2,1}(\Omega \times [0, T])$, satisfying $\langle \mathbf{n}(\mathbf{x}), \nabla h(\mathbf{x}, t) \rangle = 0$ for all $\mathbf{x} \in \partial\Omega, t \in [0, T]$. In order to prove the claim, we need to show that (A.3) holds. Throughout the proof, we will let $\lambda_n \equiv \lambda_{\delta_n}$.

Recall that, for any $n \geq 1$, ρ^{δ_n} is a weak solution of the PDE (B.2) with initial and boundary conditions (B.3). Hence, by Definition B.1, we have that

$$\begin{aligned} & \int_{\Omega^{\delta_n}} h^{\delta_n}(\mathbf{x}, T) \rho_T^{\delta_n}(\mathrm{d}\mathbf{x}) - \int_{\Omega^{\delta_n}} h^{\delta_n}(\mathbf{x}, 0) \rho_0^{\delta_n}(\mathrm{d}\mathbf{x}) \\ &= \int_0^T \int_{\Omega^{\delta_n}} \left[\partial_t h^{\delta_n}(\mathbf{x}, t) + \tau \Delta h^{\delta_n}(\mathbf{x}, t) \right. \\ & \quad \left. - \langle \nabla \Psi(\mathbf{x}, \rho_t^{\delta_n}), \nabla h^{\delta_n}(\mathbf{x}, t) \rangle \right] \rho_t^{\delta_n}(\mathrm{d}\mathbf{x}) \mathrm{d}t, \end{aligned} \quad (\text{F.14})$$

for any $h^{\delta_n} \in \mathcal{C}^{2,1}(\Omega^{\delta_n} \times [0, T])$ satisfying $\langle \mathbf{n}(\mathbf{x}), \nabla h^{\delta_n}(\mathbf{x}, t) \rangle = 0$ for all $\mathbf{x} \in \partial\Omega^{\delta_n}$, $t \in [0, T]$. Now, we set

$$h^{\delta_n}(\mathbf{x}, t) = h(\mathbf{x}/\lambda_n, t). \quad (\text{F.15})$$

By definition of Ω_n^δ , we have that $h^{\delta_n} \in \mathcal{C}^{2,1}(\Omega^{\delta_n} \times [0, T])$ since $h \in \mathcal{C}^{2,1}(\Omega \times [0, T])$. Furthermore, $\langle \mathbf{n}(\mathbf{x}), \nabla h(\mathbf{x}, t) \rangle = 0$ for all $\mathbf{x} \in \partial\Omega$, $t \in [0, T]$ immediately implies that $\langle \mathbf{n}(\mathbf{x}), \nabla h^{\delta_n}(\mathbf{x}, t) \rangle = 0$ for all $\mathbf{x} \in \partial\Omega^{\delta_n}$, $t \in [0, T]$.

Recall that

$$\Psi(\mathbf{x}, \rho_t^{\delta_n}) = V^{\delta_n}(\mathbf{x}) + U^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}) = -\nu_0 K^{\delta_n} * f(\mathbf{x}) + \nu_0 K^{\delta_n} * K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}). \quad (\text{F.16})$$

Thus, (F.14) can be rewritten as

$$\begin{aligned} & \int_{\Omega^{\delta_n}} h^{\delta_n}(\mathbf{x}, T) \rho_T^{\delta_n}(\mathrm{d}\mathbf{x}) - \int_{\Omega^{\delta_n}} h^{\delta_n}(\mathbf{x}, 0) \rho_0^{\delta_n}(\mathrm{d}\mathbf{x}) \\ &= \int_0^T \int_{\Omega^{\delta_n}} \left[\partial_t h^{\delta_n}(\mathbf{x}, t) + \nu_0 \langle \nabla K^{\delta_n} * f(\mathbf{x}), \nabla h^{\delta_n}(\mathbf{x}, t) \rangle \right] \rho_t^{\delta_n}(\mathrm{d}\mathbf{x}) \mathrm{d}t \\ &+ \int_0^T \int_{\Omega^{\delta_n}} \left[\tau \Delta h^{\delta_n}(\mathbf{x}, t) - \nu_0 \langle \nabla K^{\delta_n} * K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}), \nabla h^{\delta_n}(\mathbf{x}, t) \rangle \right] \rho_t^{\delta_n}(\mathrm{d}\mathbf{x}) \mathrm{d}t. \end{aligned} \quad (\text{F.17})$$

Since $(\rho^{\delta_n})_{n \geq 1}$ converges in $\mathcal{C}([0, T], \mathcal{P}_2(\Omega))$ to ρ^∞ by Lemma F.1, we have that

$$\begin{aligned} & \lim_{n \rightarrow \infty} \int_{\Omega^{\delta_n}} h^{\delta_n}(\mathbf{x}, T) \rho_T^{\delta_n}(\mathrm{d}\mathbf{x}) = \int_{\Omega} h(\mathbf{x}, T) \rho_T^\infty(\mathrm{d}\mathbf{x}), \\ & \lim_{n \rightarrow \infty} \int_0^T \int_{\Omega^{\delta_n}} \left[\partial_t h^{\delta_n}(\mathbf{x}, t) + \tau \Delta h^{\delta_n}(\mathbf{x}, t) \right] \rho_t^{\delta_n}(\mathrm{d}\mathbf{x}) \mathrm{d}t \\ &= \int_0^T \int_{\Omega} [\partial_t h(\mathbf{x}, t) + \tau \Delta h(\mathbf{x}, t)] \rho_t^\infty(\mathrm{d}\mathbf{x}) \mathrm{d}t. \end{aligned} \quad (\text{F.18})$$

Furthermore, since $\rho_0^{\delta_n}(\mathbf{x}) = \lambda_n^{-d} \rho_{\text{init}}(\mathbf{x}/\lambda_n)$, we have that

$$\int_{\Omega^{\delta_n}} h^{\delta_n}(\mathbf{x}, 0) \rho_0^{\delta_n}(\mathrm{d}\mathbf{x}) = \int_{\Omega} h(\mathbf{x}, 0) \rho_{\text{init}}(\mathbf{x}) \mathrm{d}\mathbf{x}. \quad (\text{F.19})$$

Let us use the notation $h_t(\mathbf{x}) = h(\mathbf{x}, t)$ and $h_t^{\delta_n}(\mathbf{x}) = h^{\delta_n}(\mathbf{x}, t)$. Again, we set $\rho_t^{\delta_n}(\mathbf{x}) = 0$ for $\mathbf{x} \notin \Omega^{\delta_n}$. We further define $\tilde{\rho}_t^{\delta_n}(\mathbf{x}) = \lambda_n^d \rho_t^{\delta_n}(\lambda_n \mathbf{x})$, which is a probability density on Ω . Since

$\rho_t^{\delta_n}(\cdot) \rightarrow \rho_t^\infty(\cdot)$ in $\mathcal{P}_2(\Omega)$ and $\lambda_n \rightarrow 1$, we have $\tilde{\rho}_t^{\delta_n}(\cdot) \rightarrow \rho_t^\infty(\cdot)$ in $\mathcal{P}_2(\Omega)$ as well. Hence

$$\lim_{n \rightarrow \infty} \int_0^T \int_{\Omega^{\delta_n}} \langle \nabla K^{\delta_n} * f(\mathbf{x}), \nabla h_t^{\delta_n}(\mathbf{x}) \rangle \rho_t^{\delta_n}(\mathbf{x}) \, d\mathbf{x} \, dt \quad (\text{F.20})$$

$$\begin{aligned} &= \lim_{n \rightarrow \infty} \lambda_n^{-1} \int_0^T \int_{\Omega} \langle \nabla K^{\delta_n} * f(\lambda_n \mathbf{x}), \nabla h_t(\mathbf{x}) \rangle \tilde{\rho}_t^{\delta_n}(\mathbf{x}) \, d\mathbf{x} \, dt \\ &= \int_0^T \int_{\Omega} \langle \nabla f(\mathbf{x}), \nabla h_t^{\delta_n}(\mathbf{x}) \rangle \rho_t^\infty(\mathbf{x}) \, d\mathbf{x} \, dt, \end{aligned} \quad (\text{F.21})$$

where the last equality follows since $\lambda_n \rightarrow 1$, and $\nabla K^{\delta_n} * f(\lambda_n \mathbf{x}) \rightarrow \nabla f(\mathbf{x})$ uniformly in Ω .

Furthermore, we have that

$$\begin{aligned} &\left| \int_0^T \int_{\Omega^{\delta_n}} \langle \nabla K^{\delta_n} * K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}), \nabla h_t^{\delta_n}(\mathbf{x}) \rangle \rho_t^{\delta_n}(\mathbf{x}) \, d\mathbf{x} \, dt \right. \\ &\quad \left. + \frac{1}{2} \int_0^T \int_{\Omega} \Delta h_t(\mathbf{x}) (\rho_t^\infty(\mathbf{x}))^2 \, d\mathbf{x} \, dt \right| \\ &\leq \left| \int_0^T \int_{\Omega^{\delta_n}} \langle \nabla K^{\delta_n} * K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}), \nabla h_t^{\delta_n}(\mathbf{x}) \rangle \rho_t^{\delta_n}(\mathbf{x}) \, d\mathbf{x} \, dt \right. \\ &\quad \left. - \int_0^T \int_{\Omega} \langle \nabla K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}), \nabla h_t(\mathbf{x}) \rangle K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}) \, d\mathbf{x} \, dt \right| \\ &+ \left| \int_0^T \int_{\Omega} \langle \nabla K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}), \nabla h_t(\mathbf{x}) \rangle K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}) \, d\mathbf{x} \, dt \right. \\ &\quad \left. + \frac{1}{2} \int_0^T \int_{\Omega} \Delta h_t(\mathbf{x}) (K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}))^2 \, d\mathbf{x} \, dt \right| \\ &+ \left| \frac{1}{2} \int_0^T \int_{\Omega} \Delta h_t(\mathbf{x}) (K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}))^2 \, d\mathbf{x} \, dt \right. \\ &\quad \left. - \frac{1}{2} \int_0^T \int_{\Omega} \Delta h_t(\mathbf{x}) (\rho_t^\infty(\mathbf{x}))^2 \, d\mathbf{x} \, dt \right|. \end{aligned} \quad (\text{F.22})$$

The second term in the right-hand side of (F.22) is equal to 0 by integration by parts. The third integral in the right-hand side of (F.22) is upper bounded as follows:

$$\begin{aligned} &\left| \frac{1}{2} \int_0^T \int_{\Omega} \Delta h_t(\mathbf{x}) (K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}))^2 \, d\mathbf{x} \, dt - \frac{1}{2} \int_0^T \int_{\Omega} \Delta h_t(\mathbf{x}) (\rho_t^\infty(\mathbf{x}))^2 \, d\mathbf{x} \, dt \right| \\ &\leq C \int_0^T \int_{\Omega} \left| (K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}))^2 - (\rho_t^\infty(\mathbf{x}))^2 \right| \, d\mathbf{x} \, dt \\ &\leq C \int_0^T \int_{\Omega} \left| K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}) - \rho_t^\infty(\mathbf{x}) \right| \left| K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}) + \rho_t^\infty(\mathbf{x}) \right| \, d\mathbf{x} \, dt \\ &\leq C \left(\int_0^T \int_{\Omega} \left| K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}) - \rho_t^\infty(\mathbf{x}) \right|^2 \, d\mathbf{x} \, dt \right)^{1/2} \\ &\quad \left(\int_0^T \int_{\Omega} \left| K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}) + \rho_t^\infty(\mathbf{x}) \right|^2 \, d\mathbf{x} \, dt \right)^{1/2}, \end{aligned} \quad (\text{F.23})$$

which converges to 0, as $(K^{\delta_n} * \rho_t^{\delta_n})_{n \geq 1}$ converges in $\mathcal{L}^2(\Omega \times [0, T])$ to ρ^∞ . The first term in the right-hand side of (F.22) is upper bounded as follows:

$$\begin{aligned}
& \left| \int_0^T \int_{\Omega^{\delta_n}} \langle \nabla K^{\delta_n} * K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}), \nabla h_t^{\delta_n}(\mathbf{x}) \rangle \rho_t^{\delta_n}(\mathbf{x}) \, d\mathbf{x} \, dt \right. \\
& \quad \left. - \int_0^T \int_{\Omega} \langle \nabla K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}), \nabla h_t(\mathbf{x}) \rangle K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}) \, d\mathbf{x} \, dt \right| \\
& \leq \left| \int_0^T \int_{\Omega^{\delta_n}} \langle \nabla K^{\delta_n} * K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}), \nabla h_t^{\delta_n}(\mathbf{x}) \rangle \rho_t^{\delta_n}(\mathbf{x}) \, d\mathbf{x} \, dt \right. \\
& \quad \left. - \int_0^T \int_{\Omega^{\delta_n}} \langle \nabla K^{\delta_n} * K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}), \nabla h_t(\mathbf{x}) \rangle \rho_t^{\delta_n}(\mathbf{x}) \, d\mathbf{x} \, dt \right| \\
& + \left| \int_0^T \int_{\Omega^{\delta_n}} \langle \nabla K^{\delta_n} * K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}), \nabla h_t(\mathbf{x}) \rangle \rho_t^{\delta_n}(\mathbf{x}) \, d\mathbf{x} \, dt \right. \\
& \quad \left. - \int_0^T \int_{\Omega} \langle \nabla K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}), \nabla h_t(\mathbf{x}) \rangle K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}) \, d\mathbf{x} \, dt \right|.
\end{aligned} \tag{F.24}$$

The first term is upper bounded using

$$\begin{aligned}
& \left| \int_0^T \int_{\Omega^{\delta_n}} \langle \nabla K^{\delta_n} * K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}), \nabla h_t^{\delta_n}(\mathbf{x}) \rangle \rho_t^{\delta_n}(\mathbf{x}) \, d\mathbf{x} \, dt \right. \\
& \quad \left. - \int_0^T \int_{\Omega} \langle \nabla K^{\delta_n} * K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}), \nabla h_t(\mathbf{x}) \rangle \rho_t^{\delta_n}(\mathbf{x}) \, d\mathbf{x} \, dt \right| \\
& \leq \| \nabla h^{\delta_n} - \nabla h \|_{\mathcal{L}^\infty(\Omega \times [0, T])} \left\| \left| \nabla K^{\delta_n} * K^{\delta_n} * \rho_t^{\delta_n} \right| \times \rho_t^{\delta_n} \right\|_{\mathcal{L}^1(\Omega^{\delta_n} \times [0, T])}.
\end{aligned} \tag{F.25}$$

Notice that

$$\begin{aligned}
& \left\| \left| \nabla K^{\delta_n} * K^{\delta_n} * \rho_t^{\delta_n} \right| \times \rho_t^{\delta_n} \right\|_{\mathcal{L}^1(\Omega^{\delta_n} \times [0, T])} \\
& \stackrel{(a)}{\leq} \left\| \left| \nabla K^{\delta_n} * K^{\delta_n} * \rho_t^{\delta_n} \right| \times \sqrt{\rho_t^{\delta_n}} \right\|_{\mathcal{L}^2(\Omega^{\delta_n} \times [0, T])} \left\| \sqrt{\rho_t^{\delta_n}} \right\|_{\mathcal{L}^2(\Omega^{\delta_n} \times [0, T])} \\
& \leq \sqrt{T} \left\| \left| \nabla K^{\delta_n} * K^{\delta_n} * \rho_t^{\delta_n} \right| \times \sqrt{\rho_t^{\delta_n}} \right\|_{\mathcal{L}^2(\Omega^{\delta_n} \times [0, T])},
\end{aligned} \tag{F.26}$$

where (a) follows from an application of Cauchy-Schwartz. By Lemma E.5, we deduce that the right-hand side of (F.26) is bounded uniformly in δ_n . Thus, the first term of (F.24) converges to 0

because of Eq. (F.25). As concerns the second term of (F.24), we have that

$$\begin{aligned}
& \left| \int_0^T \int_{\Omega^{\delta_n}} \langle \nabla K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}), \nabla h_t(\mathbf{x}) \rangle \rho_t^{\delta_n}(\mathbf{x}) \, d\mathbf{x} \, dt \right. \\
& \quad \left. - \int_0^T \int_{\Omega} \langle \nabla K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}), \nabla h_t(\mathbf{x}) \rangle K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{x}) \, d\mathbf{x} \, dt \right| \\
& \leq \left| \int_0^T \int_{\Omega} \int_{\Omega} \langle \nabla K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{y}), \nabla h_t(\mathbf{x}) \rangle K^{\delta_n}(\mathbf{x} - \mathbf{y}) \rho_t^{\delta_n}(\mathbf{x}) \, d\mathbf{y} \, d\mathbf{x} \, dt \right. \\
& \quad \left. - \int_0^T \int_{\Omega} \int_{\Omega} \langle \nabla K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{y}), \nabla h_t(\mathbf{y}) \rangle K^{\delta_n}(\mathbf{x} - \mathbf{y}) \rho_t^{\delta_n}(\mathbf{x}) \, d\mathbf{y} \, d\mathbf{x} \, dt \right| \\
& = \left| \int_0^T \int_{\Omega} \int_{\Omega} \langle \nabla K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{y}), \nabla h_t(\mathbf{x}) - \nabla h_t(\mathbf{y}) \rangle K^{\delta_n}(\mathbf{x} - \mathbf{y}) \rho_t^{\delta_n}(\mathbf{x}) \, d\mathbf{y} \, d\mathbf{x} \, dt \right| \\
& \leq C \int_0^T \int_{\Omega} \int_{\Omega} |\mathbf{x} - \mathbf{y}| |\nabla K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{y})| K^{\delta_n}(\mathbf{x} - \mathbf{y}) \rho_t^{\delta_n}(\mathbf{x}) \, d\mathbf{y} \, d\mathbf{x} \, dt.
\end{aligned} \tag{F.27}$$

Recall that $\rho_t^{\delta_n}$ is supported on $\Omega^{\delta_n} \subseteq \Omega$, and Ω is bounded. In addition, since the kernel K has bounded support, the diameter of the support of K^{δ_n} is at most δ_n times a constant. Consequently, the last term in the right-hand side of (F.27) is upper bounded by

$$\begin{aligned}
& \delta_n C_1 \int_0^T \int_{\Omega} \int_{\Omega} |\nabla K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{y})| K^{\delta_n}(\mathbf{x} - \mathbf{y}) \rho_t^{\delta_n}(\mathbf{x}) \, d\mathbf{y} \, d\mathbf{x} \, dt \\
& = \delta_n C_1 \int_0^T \int_{\Omega} |\nabla K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{y})| K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{y}) \, d\mathbf{y} \, dt \\
& \leq \delta_n C_1 \left(\int_0^T \int_{\Omega} |\nabla K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{y})|^2 \, d\mathbf{y} \, dt \right)^{1/2} \\
& \quad \left(\int_0^T \int_{\Omega} (K^{\delta_n} * \rho_t^{\delta_n}(\mathbf{y}))^2 \, d\mathbf{y} \, dt \right)^{1/2}.
\end{aligned} \tag{F.28}$$

By using that $K^{\delta_n} * \rho^{\delta_n} \in \mathcal{L}^2(\Omega \times [0, T])$ and the result of Lemma E.5, we have that the two last integrals are bounded uniformly in δ . As a result, the right-hand side of (F.28) converges to 0, which implies that the right-hand side of (F.22) also converges to 0. By putting this fact together with (F.18) and (F.21), the desired result follows. $\square$

We have now proved that $(\rho^{\delta_n})_{n \geq 1}$ converges to a weak solution of the limit PDE (A.1). In order to prove the uniqueness of the weak solutions of the limit PDE, we next prove a bound on $\|\rho_t^{\delta_n}\|_{\mathcal{L}^4(\Omega)}$, which along with Lemma A.2 proves the uniqueness claim.

Lemma F.5 (Uniform bound in $\mathcal{L}^4$). *Assume that $\rho_{\text{init}}, f \in \mathcal{C}^\infty(\Omega)$ and consider the sequence $(\rho^{\delta_n})_{n \geq 1}$. Then,*

$$\sup_{n \geq 1, t \in [0, T_0]} \|\rho_t^{\delta_n}\|_{\mathcal{L}^4(\Omega)} \leq C(\Omega)(1 + T), \tag{F.29}$$

where

$$T_0 = \frac{1}{C(\Omega)(1 + T)}, \tag{F.30}$$

for some bounded constant $0 < C(\Omega) < \infty$.

Proof. For simplicity, we indicate the norms $\mathcal{L}^p(\Omega)$ by $\|\cdot\|_p$. For a function $g \in \mathcal{C}^m(\Omega)$, we let $\nabla^{\otimes m} g$ be the vector with coordinates $\partial^m g / (\partial_{i_1} \dots \partial_{i_m})$, with $1 \leq i_1, i_2, \dots, i_m \leq d$. The proof strategy to prove this lemma is to first bound $\|\nabla^{\otimes m} \rho_t^{\delta_n}\|_2$, for some $m \geq d/4$, and then apply the Gagliardo-Nirenberg interpolation inequality (cf. Lemma H.3) to bound $\|\rho_t^{\delta_n}\|_4$. Throughout this proof, we will use C , C_k and so on to denote constants that can depend on the domain Ω , but do not depend on t or δ .

Before proceeding, we need to establish some notations and definitions.

For a function g and an integer $k \geq 0$, we denote its Sobolev norms by

$$\|g\|_{(k)} = \left(\sum_{m=0}^k \|\nabla^{\otimes m} g\|_2^2 \right)^{1/2}. \quad (\text{F.31})$$

We will use the following relations on Sobolev norms (see [Oel01, Equation (1.14)]):

$$\|g\|_{(k)} \leq \begin{cases} \bar{C}_k (\|g\|_2^2 + \|(-\Delta)^{k/2} g\|_2^2)^{1/2} & \text{if } k \text{ is even,} \\ \bar{C}_k (\|g\|_2^2 + \|\nabla(-\Delta)^{(k-1)/2} g\|_2^2)^{1/2} & \text{if } k \text{ is odd.} \end{cases} \quad (\text{F.32})$$

Instead of bounding $\|\nabla^{\otimes m} \rho_t^{\delta_n}\|_2$, we will bound the dominating quantity $\|\rho_t^{\delta_n}\|_{(m)}$. To this end, we follow a similar strategy as in [Oel01]. Namely, we derive descriptions of the evolution of $\|(-\Delta)^m \rho_t^{\delta_n}\|_2$ and $\|(-\Delta)^m (\rho_t^{\delta_n} * K^{\delta_n} - f)\|_2$. More precisely, we derive a recursive equation (on m) for the evolution of a suitably chosen linear combination of these two quantities.

Since ρ^{δ_n} is a solution of the PDE (B.2), we have

$$\partial_t (-\Delta)^m \rho_t^{\delta_n}(\mathbf{x}) = -\tau (-\Delta)^{m+1} \rho_t^{\delta_n}(\mathbf{x}) + (-\Delta)^m \nabla \cdot \left(\rho_t^{\delta_n}(\mathbf{x}) \nabla (V + U * \rho_t^{\delta_n}) \right) \quad (\text{F.33})$$

Following along the same lines as in derivation of [Oel01, Equation (3.12)], we obtain

$$\begin{aligned} & \frac{d}{dt} \|(-\Delta)^m \rho_t^{\delta_n}\|_2^2 \\ & \leq (CC_m - 2\tau) \|\nabla(-\Delta)^m \rho_t^{\delta_n}\|_2^2 \\ & \quad + \frac{2}{C} \|\rho_t^{\delta_n}\|_\infty \left\langle \nabla(-\Delta)^m (V + U * \rho_t^{\delta_n}), \rho_t^{\delta_n} \nabla(-\Delta)^m (V + U * \rho_t^{\delta_n}) \right\rangle \\ & \quad + \frac{\tilde{C}_m}{C} \sum_{q=1}^m \left(\|V + U * \rho_t^{\delta_n}\|_{(2m+1-q)}^2 \|\rho_t^{\delta_n}\|_{(q+1+d/2)}^2 \right. \\ & \quad \left. + \|V + U * \rho_t^{\delta_n}\|_{(q+1+d/2)}^2 \|\rho_t^{\delta_n}\|_{(2m+1-q)}^2 \right), \end{aligned} \quad (\text{F.34})$$

where C_m and $\tilde{C}_m$ are positive constants that depend on m and $C > 0$ is a constant which can be chosen arbitrarily.

We set $m = \lceil 1 + d/2 \rceil$ for which we can upper bound the right-hand side of (F.34) as

$$\begin{aligned}
& \frac{d}{dt} \left\| (-\Delta)^m \rho_t^{\delta_n} \right\|_2^2 \\
& \leq (CC_m - 2\tau) \left\| \nabla (-\Delta)^m \rho_t^{\delta_n} \right\|_2^2 \\
& \quad + \frac{2}{C} \|\rho_t^{\delta_n}\|_\infty \left\langle \nabla (-\Delta)^m (V + U * \rho_t), \rho_t^{\delta_n} \nabla (-\Delta)^m (V + U * \rho_t^{\delta_n}) \right\rangle \\
& \quad + \frac{2m\tilde{C}_m}{C} \left\| V + U * \rho_t^{\delta_n} \right\|_{(2m)}^2 \|\rho_t^{\delta_n}\|_{(2m)}^2. \tag{F.35}
\end{aligned}$$

We next move to the next quantity. Write

$$\begin{aligned}
& \frac{d}{dt} \left\| (-\Delta)^m (K^{\delta_n} * \rho_t^{\delta_n} - f) \right\|_2^2 \\
& \leq 2 \left\langle (-\Delta)^m (K^{\delta_n} * \rho_t^{\delta_n} - f), \partial_t (-\Delta)^m \rho_t^{\delta_n} * K^{\delta_n} \right\rangle \\
& = 2 \left\langle (-\Delta)^m (V^{\delta_n} + U^{\delta_n} * \rho_t^{\delta_n}), \partial_t (-\Delta)^m \rho_t^{\delta_n} \right\rangle \\
& = -2\tau \left\langle (-\Delta)^m (V^{\delta_n} + U^{\delta_n} * \rho_t^{\delta_n}), (-\Delta)^{m+1} \rho_t^{\delta_n} \right\rangle \\
& \quad + 2 \left\langle (-\Delta)^m (V^{\delta_n} + U^{\delta_n} * \rho_t^{\delta_n}), (-\Delta)^m \nabla \cdot \left(\rho_t^{\delta_n}(\mathbf{x}) \nabla (V^{\delta_n} + U^{\delta_n} * \rho_t^{\delta_n}) \right) \right\rangle, \tag{F.36}
\end{aligned}$$

where the last step follows from (F.33). Note that the first term on the right-hand side can be bounded as

$$\begin{aligned}
& -2\tau \left\langle (-\Delta)^m (V^{\delta_n} + U^{\delta_n} * \rho_t^{\delta_n}), (-\Delta)^{m+1} \rho_t^{\delta_n} \right\rangle \\
& = -2\tau \left\langle (-\Delta)^{m+1} (K^{\delta_n} * f), (-\Delta)^m \rho_t^{\delta_n} \right\rangle - 2\tau \left\| \nabla (-\Delta)^m (K^{\delta_n} * \rho_t^{\delta_n}) \right\|_2^2 \\
& \leq 2\tau \|(-\Delta)^{m+1} f\|_2 \|(-\Delta)^m \rho_t^{\delta_n}\|_2 - 2\tau \left\| \nabla (-\Delta)^m (K^{\delta_n} * \rho_t^{\delta_n}) \right\|_2^2, \tag{F.37}
\end{aligned}$$

where the last step follows from Young's convolution inequality and the fact that $\|K^{\delta_n}\|_1 = 1$.

The second term in (F.36) can be bounded following the same lines as in derivation of [Oel01, Equations (3.3) and (3.16)], which along with (F.37) gives

$$\begin{aligned}
& \frac{d}{dt} \left\| (-\Delta)^m (K^{\delta_n} * \rho_t^{\delta_n} - f) \right\|_2^2 \\
& \leq CC_m \left\| \nabla (-\Delta)^m (K^{\delta_n} * \rho_t^{\delta_n} - f) \right\|_2^2 + 2\tau \|(-\Delta)^{m+1} f\|_2 \|(-\Delta)^m \rho_t^{\delta_n}\|_2 \\
& \quad - 2\tau \left\| \nabla (-\Delta)^m (K^{\delta_n} * \rho_t^{\delta_n}) \right\|_2^2 \\
& \quad - 2 \left\langle \nabla (-\Delta)^m (V + U * \rho_t^{\delta_n}), \rho_t^{\delta_n} \nabla (-\Delta)^m (V + U * \rho_t^{\delta_n}) \right\rangle \\
& \quad + \frac{\tilde{C}_m}{C} \sum_{q=1}^m \left(\left\| V + U * \rho_t^{\delta_n} \right\|_{(2m+1-q)}^2 \|\rho_t^{\delta_n}\|_{(q+1+d/2)}^2 \right. \\
& \quad \left. + \left\| V + U * \rho_t^{\delta_n} \right\|_{(q+1+d/2)}^2 \|\rho_t^{\delta_n}\|_{(2m+1-q)}^2 \right). \tag{F.38}
\end{aligned}$$

Since $f \in \mathcal{C}^\infty(\Omega)$, there exists constant $M > 0$, such that $\|(-\Delta)^{m+1}f\|_2 \leq M$, $\|\nabla(-\Delta)^m f\|_2 \leq M$. Using the particular choice of m , we can upper bound the right-hand side of (F.38) as

$$\begin{aligned}
& \frac{d}{dt} \left\| (-\Delta)^m (K^{\delta_n} * \rho_t^{\delta_n} - f) \right\|_2^2 \\
& \leq (2CC_m - 2\tau) \left\| \nabla(-\Delta)^m (K^{\delta_n} * \rho_t^{\delta_n}) \right\|_2^2 + 2\tau M \left\| (-\Delta)^m \rho_t^{\delta_n} \right\|_2 + 2M^2 CC_m \\
& \quad - 2 \left\langle \nabla(-\Delta)^m (V + U * \rho_t^{\delta_n}), \rho_t^{\delta_n} \nabla(-\Delta)^m (V + U * \rho_t^{\delta_n}) \right\rangle \\
& \quad + \frac{2m\tilde{C}_m}{C} \left\| V + U * \rho_t^{\delta_n} \right\|_{(2m)}^2 \left\| \rho_t^{\delta_n} \right\|_{(2m)}^2.
\end{aligned} \tag{F.39}$$

Define $C_1 \equiv 2\|\rho_{\text{init}}\|_{(2m)}$ and let

$$T_n \equiv \inf \left\{ t \in [0, T] : \left\| \rho_t^{\delta_n} \right\|_{(2m)} > C_1 \right\} \wedge T,$$

for $n \geq 1$. Clearly, $T_n > 0$ by choice of C_1 . In addition, by applying Sobolev's inequality (see e.g. [Oel01, Equation (1.12)]), we have

$$\left\| \rho_t^{\delta_n} \right\|_\infty \leq C_2 \left\| \rho_t^{\delta_n} \right\|_{(2m)} \leq C_1 C_2, \quad \text{for } t \in [0, T_n], n \geq 1.$$

where $C_2 > 0$ is a constant depending on d . We let $C_* \equiv C_1 C_2 / C$. Recall that the constant $C > 0$ in (F.35) and (F.39) was arbitrary. We choose it in a way that $C < \tau / (2C_m)$. We then consider the evolution of the following linear combination of the two quantities we analyzed above. Note that by Equations (F.35) and (F.39), we have for $t \in [0, T_n]$,

$$\begin{aligned}
& \frac{d}{dt} \left(\left\| (-\Delta)^m \rho_t^{\delta_n} \right\|_2^2 + C_* \left\| (-\Delta)^m (K^{\delta_n} * \rho_t^{\delta_n} - f) \right\|_2^2 \right) \\
& \leq -\frac{3\tau}{2} \left\| \nabla(-\Delta)^m \rho_t^{\delta_n} \right\|_2^2 - C_* \tau \left\| \nabla(-\Delta)^m (K^{\delta_n} * \rho_t^{\delta_n}) \right\|_2^2 \\
& \quad + C_* \left(2\tau M \left\| (-\Delta)^m \rho_t^{\delta_n} \right\|_2 + \tau M^2 \right) \\
& \quad + \frac{2m\tilde{C}_m}{C} (1 + C_*) \left\| V + U * \rho_t^{\delta_n} \right\|_{(2m)}^2 \left\| \rho_t^{\delta_n} \right\|_{(2m)}^2 \\
& \leq C_* \left(2\tau M \left\| (-\Delta)^m \rho_t^{\delta_n} \right\|_2 + \tau M^2 \right) \\
& \quad + \frac{2m\tilde{C}_m}{C} (1 + C_*) \left(\left\| \rho_t^{\delta_n} \right\|_{(2m)}^2 + \left\| V + U * \rho_t^{\delta_n} \right\|_{(2m)}^2 \right)^2 \\
& \stackrel{(a)}{\leq} \tau M^2 C_* + C_3 \left(\left\| \rho_t^{\delta_n} \right\|_{(2m)}^2 + C_* \left\| V + U * \rho_t^{\delta_n} \right\|_{(2m)}^2 \right)^2 \\
& \stackrel{(b)}{\leq} \tau M^2 C_* + C_3 \left(\left\| \rho_t^{\delta_n} \right\|_{(2m)}^2 + C_* \left\| -f + K^{\delta_n} * \rho_t^{\delta_n} \right\|_{(2m)}^2 \right)^2,
\end{aligned} \tag{F.40}$$

where in (a) we use the fact that $\left\| (-\Delta)^m \rho_t^{\delta_n} \right\|_2 \leq \left\| \rho_t^{\delta_n} \right\|_{(2m)}$, which follows immediately from (F.31); (b) follows from the fact that for any function $g \in \mathcal{L}^2(\Omega)$, $\|g * K^{\delta_n}\|_2 \leq \|K^{\delta_n}\|_1 \|g\|_2 = \|g\|_2$, by Young's inequality for convolution.

Another observation that will be used later is that

$$\frac{d}{dt} \left(\left\| \rho_t^{\delta_n} \right\|_2^2 + C_* \left\| K^{\delta_n} * \rho_t^{\delta_n} - f \right\|_2^2 \right) \leq 0. \tag{F.41}$$

This claim follows by repeating the same argument we had to derive (F.40), for $m = 0$. In this case, we have analogous equations to (F.35) and (F.39), where only the first two terms appear.

Next note that by (F.32), we have for $t \in [0, T_n]$,

$$\begin{aligned}
& \|\rho_t^{\delta_n}\|_{(2m)}^2 + C_* \|K^{\delta_n} * \rho_t^{\delta_n} - f\|_{(2m)}^2 \\
& \leq \bar{C}_m \left(\|\rho_t^{\delta_n}\|_2^2 + \|(-\Delta)^m \rho_t^{\delta_n}\|_2^2 \right. \\
& \quad \left. + C_* \left(\|K^{\delta_n} * \rho_t^{\delta_n} - f\|_2^2 + \|(-\Delta)^m (K^{\delta_n} * \rho_t^{\delta_n} - f)\|_2^2 \right) \right) \\
& \leq \bar{C}_m \left(\|\rho_{\text{init}}\|_{(2m)}^2 + C_* \|K^{\delta_n} * \rho_{\text{init}} - f\|_{(2m)}^2 \right) \\
& \quad + \bar{C}_m \tau M^2 C_* t + \bar{C}_m C_3 \int_0^t \left(\|\rho_s^{\delta_n}\|_{(2m)}^2 + C_* \|K^{\delta_n} * \rho_s^{\delta_n} - f\|_{(2m)}^2 \right) ds, \tag{F.42}
\end{aligned}$$

where the last step is a result of (F.41) and (F.40). Let us stress that $\bar{C}_m$, C_* , C_3 are constants that are independent of n .

We further note that

$$\begin{aligned}
\|\rho_{\text{init}}\|_{(2m)}^2 + C_* \|K^{\delta_n} * \rho_{\text{init}} - f\|_{(2m)}^2 & \leq \|\rho_{\text{init}}\|_{(2m)}^2 (1 + 2C_*) + 2C_* \|f\|_{(2m)}^2 \\
& \leq (1 + 2C_*) \frac{C_1^2}{4} + 2C_* \|f\|_{(2m)}^2, \tag{F.43}
\end{aligned}$$

for $n \geq 1$. Here, the first step is a result of triangle inequality and the Young's inequality for convolution along with the fact that $\|K^{\delta_n}\|_1 = 1$. The second step follows from definition of C_1 . Since $f \in \mathcal{C}^\infty(\Omega)$, $\|f\|_{(2m)}^2$ is uniformly bounded over Ω . We denote the right-hand side of (F.43) by the constant C_4 . Using bound (F.43) into (F.42) results in

$$\begin{aligned}
& \|\rho_t^{\delta_n}\|_{(2m)}^2 + C_* \|K^{\delta_n} * \rho_t^{\delta_n} - f\|_{(2m)}^2 \\
& \leq \bar{C}_m C_4 + \bar{C}_m \tau M^2 C_* t \\
& \quad + \bar{C}_m C_3 \int_0^t \left(\|\rho_s^{\delta_n}\|_{(2m)}^2 + C_* \|K^{\delta_n} * \rho_s^{\delta_n} - f\|_{(2m)}^2 \right) ds, \tag{F.44}
\end{aligned}$$

for $t \in [0, T_n]$. By employing a generalization of Gronwall's inequality (cf. Lemma H.2 and Remark H.1) we get

$$\|\rho_t^{\delta_n}\|_{(2m)}^2 + C_* \|K^{\delta_n} * \rho_t^{\delta_n} - f\|_{(2m)}^2 \leq \frac{\bar{C}_m C_4 + \bar{C}_m \tau M^2 C_* T_n}{1 - (\bar{C}_m C_4 + \bar{C}_m \tau M^2 C_* T_n) \bar{C}_m C_3 t}. \tag{F.45}$$

Therefore, for $t \in [0, T_0]$, with

$$T_0 = \frac{1}{2\bar{C}_m C_4 + 2\bar{C}_m \tau M^2 C_* T}, \tag{F.46}$$

we have that

$$\|\rho_t^{\delta_n}\|_{(2m)}^2 + C_* \|K^{\delta_n} * \rho_t^{\delta_n} - f\|_{(2m)}^2 \leq C_5 + C_6 T, \tag{F.47}$$

with $C_5 \equiv \bar{C}_m C_4$ and $C_6 \equiv \bar{C}_m \tau M^2 C_*$. Note that C_5 , C_6 and T_0 are independent of n , but depend on d . Let $m_0 = \lceil d/4 \rceil$. Then, by the choice of $m = 1 + \lceil d/2 \rceil$ we have $\|\nabla^{\otimes m_0} \rho_t^{\delta_n}\|_2 \leq \|\rho_t^{\delta_n}\|_{(2m)}$, and hence as a result of (F.47), we obtain

$$\sup_{n \geq 1, t \in [0, T_0]} \|\nabla^{\otimes m_0} \rho_t^{\delta_n}\|_2 \leq C_5 + C_6 T, \quad (\text{F.48})$$

Finally, by applying Gagliardo-Nirenberg interpolation inequality (cf. Lemma H.3) we get

$$\sup_{n \geq 1, t \in [0, T_0]} \|\rho_t^{\delta_n}\|_4 \leq C_7 + C_8 T,$$

for some constant $C_7, C_8 > 0$, which completes the proof. $\square$

Lemma F.6 (Convergence to the unique weak solution of limit PDE). *Let ρ^∞ be the limit in $\mathcal{C}([0, T], \mathcal{P}_2(\Omega))$ of the converging sequence $(\rho^{\delta_n})_{n \geq 1}$. Then, ρ^∞ is the unique weak solution of the PDE (A.1) in $\mathcal{L}^4(\Omega \times [0, T])$ with initial and boundary conditions (A.2).*

Proof. From Lemma F.3, we have that the sequence $(K^{\delta_n} * \rho^{\delta_n})_{n \geq 1}$ converges in $\mathcal{L}^2(\Omega \times [0, T])$ to ρ^∞ . Furthermore, by Lemma F.5, $\|\rho_t^\delta\|_{\mathcal{L}^4(\Omega)} \leq C(1 + T)$ for any $t \in [0, T_0]$, where C is a universal constant. By using Young's convolution inequality, we also deduce that $\|K^{\delta_n} * \rho_t^\delta\|_{\mathcal{L}^4(\Omega)} \leq C(1 + T)$ for any $t \in [0, T_0]$.

Note that $\mathcal{L}^4(\Omega)$ is a reflexive Banach space. Thus, by applying the Banach-Alaoglu theorem, every bounded sequence in $\mathcal{L}^4(\Omega)$ has a weakly convergent subsequence. This means that there exist a subsequence $K^{\delta_{n_k}} * \rho^{\delta_{n_k}}$ and a function $\tilde{\rho} \in \mathcal{L}^4(\Omega)$ such that, for any $g \in \mathcal{L}^{4/3}(\Omega)$, we have

$$\int_{\Omega} (K^{\delta_{n_k}} * \rho^{\delta_{n_k}}(\mathbf{x}) - \tilde{\rho}(\mathbf{x}))g(\mathbf{x}) d\mathbf{x} \rightarrow 0. \quad (\text{F.49})$$

Now, since Ω is bounded, $K^{\delta_{n_k}} * \rho^{\delta_{n_k}}$ and $\tilde{\rho}$ are also in $\mathcal{L}^2(\Omega)$ (as they are in $\mathcal{L}^4(\Omega)$). Thus, $K^{\delta_{n_k}} * \rho^{\delta_{n_k}} - \tilde{\rho}$ is in $\mathcal{L}^2(\Omega)$, hence it is also in $\mathcal{L}^{4/3}(\Omega)$. As a result, we can pick $g = K^{\delta_{n_k}} * \rho^{\delta_{n_k}} - \tilde{\rho}$ and obtain

$$\int_{\Omega} |K^{\delta_{n_k}} * \rho^{\delta_{n_k}}(\mathbf{x}) - \tilde{\rho}(\mathbf{x})|^2 d\mathbf{x} \rightarrow 0. \quad (\text{F.50})$$

Therefore, $\tilde{\rho}$ is the limit in $\mathcal{L}^2(\Omega)$ of the sequence $K^{\delta_{n_k}} * \rho^{\delta_{n_k}}$. By uniqueness of the limit, we conclude that $\tilde{\rho} = \rho^\infty$. As a result, $\rho^\infty \in \mathcal{L}^4(\Omega)$ for any $t \in [0, T_0]$, which implies that $\rho^\infty \in \mathcal{L}^4(\Omega \times [0, T_0])$. Thus, by Lemma F.4 and Lemma A.2, ρ^∞ is the unique weak solution of the PDE (A.1) for $t \in [0, T_0]$. Note that T_0 is decreasing with T . Thus, we can repeat the same argument with $T - T_0$ instead of T and obtain that ρ^∞ is the unique weak solution of the PDE (A.1) for $t \in [T_0, 2T_0]$. By iterating this procedure T/T_0 times, the result follows. $\square$

At this point, we state and prove a lemma showing that the sequence $(\rho_t^{\delta_n})_{n \geq 1}$ converges in $\mathcal{L}^2(\Omega)$ to ρ_t^∞ .

Lemma F.7. *For almost all $t \in [0, T]$, the measure ρ_t^∞ is the limit in $\mathcal{L}^2(\Omega)$ of the sequence $(\rho_t^{\delta_n})_{n \geq 1}$.*

Proof. The proof is similar to that of Lemma F.3. Suppose that $t \in [0, T_0]$, where T_0 is defined in the statement of Lemma F.5. Note that, for any $n \geq 1$, $\rho^{\delta_n} \in \mathcal{L}^2(\Omega)$. Let us show that $(\rho^{\delta_n})_{n \geq 1}$ is a Cauchy sequence in $\mathcal{L}^2(\Omega)$.

As $\rho_t^{\delta_n} \in \mathcal{L}^2(\Omega)$ for every $t \in [0, T_0]$, its Fourier transform exists and we denote it by $\widehat{\rho_t^{\delta_n}}$. Hence, by applying Parseval's theorem, we have

$$\begin{aligned} \limsup_{n, n' \rightarrow \infty} \|\rho^{\delta_n} - \rho^{\delta_{n'}}\|_{\mathcal{L}^2(\Omega)}^2 &= \limsup_{n, n' \rightarrow \infty} \|\rho_t^{\delta_n} - \rho_t^{\delta_{n'}}\|_{\mathcal{L}^2(\Omega)}^2 \\ &= \limsup_{n, n' \rightarrow \infty} \int_{\mathbb{R}^d} |\widehat{\rho_t^{\delta_n}}(\boldsymbol{\lambda}) - \widehat{\rho_t^{\delta_{n'}}}(\boldsymbol{\lambda})|^2 d\boldsymbol{\lambda}. \end{aligned} \quad (\text{F.51})$$

Fix $\Lambda > 1$ and decompose the integral in the right-hand side of (F.51) as

$$\limsup_{n, n' \rightarrow \infty} \int_{|\boldsymbol{\lambda}| < \Lambda} |\widehat{\rho_t^{\delta_n}}(\boldsymbol{\lambda}) - \widehat{\rho_t^{\delta_{n'}}}(\boldsymbol{\lambda})|^2 d\boldsymbol{\lambda} + \limsup_{n, n' \rightarrow \infty} \int_{|\boldsymbol{\lambda}| \geq \Lambda} |\widehat{\rho_t^{\delta_n}}(\boldsymbol{\lambda}) - \widehat{\rho_t^{\delta_{n'}}}(\boldsymbol{\lambda})|^2 d\boldsymbol{\lambda}. \quad (\text{F.52})$$

Consider the first term of (F.52). By Lemma F.1, and since by Jensen's inequality $W_1(\rho_1, \rho_2) \leq W_2(\rho_1, \rho_2)$ for any two distributions ρ_1, ρ_2 , we have $W_1(\rho_t^{\delta_n} - \rho_t^{\delta_{n'}}) \rightarrow 0$, as $n, n' \rightarrow \infty$. Since for the complex exponential functions $\|e^{i\langle \boldsymbol{\lambda}, \boldsymbol{x} \rangle}\|_{\text{Lip}} \leq |\boldsymbol{\lambda}|$, by definition of 1-Wasserstein distance, the integrand in the first term converges pointwise to 0. Furthermore, the integrand is upper bounded by an integrable function, since $|\widehat{\rho_t^{\delta_n}}(\boldsymbol{\lambda})| \leq \|\rho_t^{\delta_n}\|_{\mathcal{L}^2(\Omega)} \leq C$ for all n and every $t \in [0, T_0]$. Hence, by dominated convergence, the first integral in (F.52) converges to 0.

As for the second term of (F.52), the following chain of inequalities holds:

$$\begin{aligned} \limsup_{n, n' \rightarrow \infty} \int_{|\boldsymbol{\lambda}| \geq \Lambda} |\widehat{\rho_t^{\delta_n}}(\boldsymbol{\lambda}) - \widehat{\rho_t^{\delta_{n'}}}(\boldsymbol{\lambda})|^2 d\boldsymbol{\lambda} &\leq 4 \sup_{n \geq 1} \int_{|\boldsymbol{\lambda}| \geq \Lambda} |\widehat{\rho_t^{\delta_n}}(\boldsymbol{\lambda})|^2 d\boldsymbol{\lambda} \\ &\leq \frac{4}{\Lambda^2} \sup_{n \geq 1} \int_{|\boldsymbol{\lambda}| \geq \Lambda} |\boldsymbol{\lambda}|^2 |\widehat{\rho_t^{\delta_n}}(\boldsymbol{\lambda})|^2 d\boldsymbol{\lambda} \\ &\leq \frac{4}{\Lambda^2} \sup_{n \geq 1} \int_{\mathbb{R}^d} |\boldsymbol{\lambda}|^2 |\widehat{\rho_t^{\delta_n}}(\boldsymbol{\lambda})|^2 d\boldsymbol{\lambda} \\ &= \frac{4}{\Lambda^2} \sup_{n \geq 1} \int_{\mathbb{R}^d} |\nabla \widehat{\rho_t^{\delta_n}}(\boldsymbol{\lambda})|^2 d\boldsymbol{\lambda} \\ &= \frac{4}{\Lambda^2} \sup_{n \geq 1} \int_{\Omega} |\nabla \rho_t^{\delta_n}(\boldsymbol{x})|^2 d\boldsymbol{x}, \end{aligned} \quad (\text{F.53})$$

where in the last equality we have applied again Parseval's theorem. In the proof of Lemma F.5, we provide an upper bound, which does not depend on n , on the Sobolev norm of $\rho_t^{\delta_n}$ (see (F.47)). Thus, as $\Lambda \rightarrow \infty$, the second term of (F.52) converges to 0.

By iterating the argument T/T_0 times, we obtain that $(\rho_t^{\delta_n})_{n \geq 1}$ is a Cauchy sequence in $\mathcal{L}^2(\Omega)$ for $t \in [0, T]$. Let $\tilde{\rho}_t^\infty \in \mathcal{L}^2(\Omega)$ be its limit. Furthermore, by Lemma F.1, $(\rho_t^{\delta_n})_{n \geq 1}$ has limit ρ_t^∞ in $\mathcal{C}([0, T], \mathcal{P}_2(\Omega))$. Therefore, the measure ρ_t^∞ has for almost every $t \in [0, T]$ the density $\tilde{\rho}_t^\infty \in \mathcal{L}^2(\Omega)$, and the proof is complete. $\square$

Theorem 5.2 follows from Lemma A.2, Lemma F.6 and Lemma F.7.

Let us define the free energy associated to the PDE (A.1) as

$$\begin{aligned} F(\rho) &= \frac{1}{2} R(\rho) - \tau S(\rho) \\ &= \frac{\nu_0}{2} \|f - \rho\|_{\mathcal{L}^2(\Omega)}^2 + \tau \int \rho(\mathbf{x}, t) \log \rho(\mathbf{x}, t) d\mathbf{x}. \end{aligned} \quad (\text{F.54})$$

As explained in Section 3.5, this limit free energy is displacement convex, and hence its W_2 gradient flow converges to the unique minimizer of (F.54). These facts are stated and proved formally in the theorem that follows.

Theorem F.8. *Assume that the initial condition $\rho^\infty(0) \in \mathcal{C}^\infty(\Omega)$. Then, the following results hold:*

1. *There exists a unique minimizer in $\mathcal{P}_2(\Omega)$, call it ρ^* , of the free energy F defined in (F.54).*
2. *For any $t \geq 0$, we have*

$$F(\rho^\infty(t)) - F(\rho^*) \leq (F(\rho^\infty(0)) - F(\rho^*))e^{-2\alpha t}, \quad (\text{F.55})$$

where α is defined in (3.1).

3. *For any $n \geq 1$ and for almost any $t \geq 0$, we have*

$$F(\rho^\delta(t)) - F(\rho^*) \leq (F(\rho^\infty(0)) - F(\rho^*))e^{-2\alpha t} + \Delta(\delta, T, d), \quad (\text{F.56})$$

where α is defined in (3.1) and $\Delta(\delta, T, d) \rightarrow 0$ as $\delta \rightarrow 0$.

Proof. The proof follows from the results of [CJM⁺01]. The technical assumptions required by [CJM⁺01] are satisfied by the PDE (A.1), since Ω is convex and bounded, the initial condition $\rho^\infty(0) \in \mathcal{L}^\infty(\Omega)$, and f satisfies the assumptions (A2) and (A3). Note also that the condition $\inf_\Omega V = 0$ coming from assumption (HV3) of [CJM⁺01] can be relaxed. In fact, adding a constant to V does not change the entropy functional in [CJM⁺01, Eq. (3)] (which corresponds to the free energy (F.54)) and the PDE in [CJM⁺01, Eq. (46)] (which corresponds to the PDE (A.1)).

The uniqueness of the minimizer ρ^* follows from [CJM⁺01, Lemma 6], which proves the first result. Since ρ^∞ is the unique weak solution of the PDE (A.1) with initial and boundary conditions (A.2), then it coincides with the unique, non-negative mass-preserving solution of [CJM⁺01, Theorem 16]. Thus, the inequality (F.55) readily follows from [CJM⁺01, Theorem 16].

It remains to prove inequality (F.56). By definition of free energy, we obtain

$$F(\rho^\delta(t)) - F(\rho^\infty(t)) = \frac{1}{2}(R(\rho^\delta(t)) - R(\rho^\infty(t))) - \tau(S(\rho^\delta(t)) - S(\rho^\infty(t))). \quad (\text{F.57})$$

Recall that, by Lemma F.7, $\rho^\delta(t)$ converges to $\rho^\infty(t)$ in $\mathcal{L}^2(\Omega)$. Consequently, by using the triangle inequality, we have that the term $R(\rho^\delta(t)) - R(\rho^\infty(t))$ tends to 0 as $\delta \rightarrow 0$.

In order to complete the proof, it remains to show that $S(\rho^\delta(t)) - S(\rho^\infty(t))$ tends to 0 as $\delta \rightarrow 0$. To do so, define

$$\begin{aligned} A &= \{\mathbf{x} \in \Omega : |\rho^\delta(\mathbf{x}, t) - \rho^\infty(\mathbf{x}, t)| > 1/4\}, \\ B &= \{\mathbf{x} \in \Omega : \rho^\delta(\mathbf{x}, t) \in [0, 1/2], \rho^\infty(\mathbf{x}, t) \in [0, 1/2]\}, \\ C &= \{\mathbf{x} \in \Omega : \rho^\delta(\mathbf{x}, t) > 1/4, \rho^\infty(\mathbf{x}, t) > 1/4\}. \end{aligned} \quad (\text{F.58})$$

Note that $A \cup B \cup C = \Omega$. In fact, suppose that $\mathbf{x} \notin B$ and $\mathbf{x} \notin C$. Then, one between $\rho^\delta(\mathbf{x}, t)$ and $\rho^\infty(\mathbf{x}, t)$ is $\in [0, 1/4]$ and the other is $> 1/2$. Consequently, $|\rho^\delta(\mathbf{x}, t) - \rho^\infty(\mathbf{x}, t)| > 1/4$ and $\mathbf{x} \in A$. This immediately implies that

$$\begin{aligned} |S(\rho^\delta(t)) - S(\rho^\infty(t))| &\leq \left| \int_A \left(\rho^\delta(\mathbf{x}, t) \log \rho^\delta(\mathbf{x}, t) - \rho^\infty(\mathbf{x}, t) \log \rho^\infty(\mathbf{x}, t) \right) d\mathbf{x} \right| \\ &\quad + \left| \int_B \left(\rho^\delta(\mathbf{x}, t) \log \rho^\delta(\mathbf{x}, t) - \rho^\infty(\mathbf{x}, t) \log \rho^\infty(\mathbf{x}, t) \right) d\mathbf{x} \right| \\ &\quad + \left| \int_C \left(\rho^\delta(\mathbf{x}, t) \log \rho^\delta(\mathbf{x}, t) - \rho^\infty(\mathbf{x}, t) \log \rho^\infty(\mathbf{x}, t) \right) d\mathbf{x} \right|. \end{aligned} \quad (\text{F.59})$$

We will now upper bound the three integrals in the RHS of (F.59). As for the first term, note that

$$\int_\Omega |\rho^\delta(\mathbf{x}, t) - \rho^\infty(\mathbf{x}, t)|^2 d\mathbf{x} \geq \int_A |\rho^\delta(\mathbf{x}, t) - \rho^\infty(\mathbf{x}, t)|^2 d\mathbf{x} \geq \frac{|A|}{16}, \quad (\text{F.60})$$

where $|A|$ denotes the volume of A . Furthermore,

$$\begin{aligned} &\left| \int_A \left(\rho^\delta(\mathbf{x}, t) \log \rho^\delta(\mathbf{x}, t) - \rho^\infty(\mathbf{x}, t) \log \rho^\infty(\mathbf{x}, t) \right) d\mathbf{x} \right| \\ &\leq \left| \int_A \rho^\delta(\mathbf{x}, t) \log \rho^\delta(\mathbf{x}, t) d\mathbf{x} \right| + \left| \int_A \rho^\infty(\mathbf{x}, t) \log \rho^\infty(\mathbf{x}, t) d\mathbf{x} \right| \\ &\leq |A|^{1/2} \left(\int_A \left(\rho^\delta(\mathbf{x}, t) \log \rho^\delta(\mathbf{x}, t) \right)^2 d\mathbf{x} \right)^{1/2} \\ &\quad + |A|^{1/2} \left(\int_A \left(\rho^\infty(\mathbf{x}, t) \log \rho^\infty(\mathbf{x}, t) \right)^2 d\mathbf{x} \right)^{1/2}. \end{aligned} \quad (\text{F.61})$$

Note that $|t \log t| \leq 1$ for $t \in [0, 1]$ and $|\log t| \leq t$ for $t \geq 1$. Thus, the RHS of (F.61) is upper bounded by

$$|A|^{1/2} \left(|A| + \|\rho^\delta(t)\|_{\mathcal{L}^4(\Omega)} \right)^{1/2} + |A|^{1/2} \left(|A| + \|\rho^\infty(t)\|_{\mathcal{L}^4(\Omega)} \right)^{1/2}. \quad (\text{F.62})$$

By Lemma F.7, for almost all $t \in [0, T]$, $\rho^\delta(t)$ converges to $\rho^\infty(t)$ in $\mathcal{L}^2(\Omega)$. Thus, by (F.60), $|A|$ tends to 0 as $\delta \rightarrow 0$. By Lemma F.6, $\rho^\infty(t) \in \mathcal{L}^4(\Omega)$ for almost all $t \in [0, T]$. Furthermore, by Lemma F.5, the quantity $\|\rho^\delta(t)\|_{\mathcal{L}^4(\Omega)}$ has a δ -free upper bound for $t \in [0, T_0]$. As a result, for almost all $t \in [0, T_0]$, the first integral in (F.59) tends to 0 as $\delta \rightarrow 0$. By iterating this argument T/T_0 times, we conclude that for almost all $t \in [0, T_0]$, the first integral in (F.59) tends to 0 as $\delta \rightarrow 0$.

In order to bound the second integral in (F.59), we write

$$\begin{aligned} &\left| \int_B \left(\rho^\delta(\mathbf{x}, t) \log \rho^\delta(\mathbf{x}, t) - \rho^\infty(\mathbf{x}, t) \log \rho^\infty(\mathbf{x}, t) \right) d\mathbf{x} \right| \\ &\leq \int_B \left| \rho^\delta(\mathbf{x}, t) \log \rho^\delta(\mathbf{x}, t) - \rho^\infty(\mathbf{x}, t) \log \rho^\infty(\mathbf{x}, t) \right| d\mathbf{x} \\ &\leq \int_B \left| \rho^\delta(\mathbf{x}, t) - \rho^\infty(\mathbf{x}, t) \right| \log \frac{1}{|\rho^\delta(\mathbf{x}, t) - \rho^\infty(\mathbf{x}, t)|} d\mathbf{x}, \end{aligned} \quad (\text{F.63})$$

where in the last inequality we have applied [CT06, Theorem 17.3.3], since $\rho^\delta(\mathbf{x}, t), \rho^\infty(\mathbf{x}, t) \in [0, 1/2]$ by definition of B . Note that

$$|\log t| \leq \max \left(2\sqrt{t}, \frac{1}{t} \right) \leq 2\sqrt{t} + \frac{1}{t}.$$

Thus, the RHS of (F.63) is upper bounded by

$$\begin{aligned} & \int_B \left| \rho^\delta(\mathbf{x}, t) - \rho^\infty(\mathbf{x}, t) \right|^2 d\mathbf{x} + 2 \int_B \left| \rho^\delta(\mathbf{x}, t) - \rho^\infty(\mathbf{x}, t) \right|^{1/2} d\mathbf{x} \\ & \|\rho^\delta(t) - \rho^\infty(t)\|_{\mathcal{L}^2(\Omega)}^2 + 2|\Omega|^{1/2} \|\rho^\delta(t) - \rho^\infty(t)\|_{\mathcal{L}^1(\Omega)}^{1/2}, \end{aligned} \quad (\text{F.64})$$

where in the last step we have used Cauchy-Schwarz inequality. By Lemma F.7, for almost all $t \in [0, T]$, $\rho^\delta(t)$ converges to $\rho^\infty(t)$ in $\mathcal{L}^2(\Omega)$. As a result, the second integral in (F.59) also tends to 0 as $\delta \rightarrow 0$.

Finally, let us bound the third integral in (F.59). Define $h(x) = x \log x$. Then, for $x > 1/4$,

$$|h'(x)| \leq 1 + |\log x| \leq 1 + \log 4 + x. \quad (\text{F.65})$$

Thus,

$$\begin{aligned} & \left| \int_C \left(\rho^\delta(\mathbf{x}, t) \log \rho^\delta(\mathbf{x}, t) - \rho^\infty(\mathbf{x}, t) \log \rho^\infty(\mathbf{x}, t) \right) d\mathbf{x} \right| \\ & \leq \int_C \left| \rho^\delta(\mathbf{x}, t) \log \rho^\delta(\mathbf{x}, t) - \rho^\infty(\mathbf{x}, t) \log \rho^\infty(\mathbf{x}, t) \right| d\mathbf{x} \\ & \leq \int_C \left| \rho^\delta(\mathbf{x}, t) - \rho^\infty(\mathbf{x}, t) \right| \cdot (1 + \log 4 + \rho^\delta(\mathbf{x}, t) + \rho^\infty(\mathbf{x}, t)) d\mathbf{x} \\ & \leq (1 + \log 4) \|\rho^\delta(t) - \rho^\infty(t)\|_{\mathcal{L}^1(\Omega)} + \|(\rho^\delta(t))^2 - (\rho^\infty(t))^2\|_{\mathcal{L}^1(\Omega)} \\ & \leq (1 + \log 4) \|\rho^\delta(t) - \rho^\infty(t)\|_{\mathcal{L}^1(\Omega)} \\ & \quad + \|\rho^\delta(t) - \rho^\infty(t)\|_{\mathcal{L}^2(\Omega)} \cdot \|\rho^\delta(t) + \rho^\infty(t)\|_{\mathcal{L}^2(\Omega)}, \end{aligned} \quad (\text{F.66})$$

where in the last step we have used Cauchy-Schwarz inequality. By Lemma F.7, for almost all $t \in [0, T]$, $\rho^\delta(t)$ converges to $\rho^\infty(t)$ in $\mathcal{L}^2(\Omega)$. By Lemma F.6, $\rho^\infty(t) \in \mathcal{L}^2(\Omega)$ for almost all $t \in [0, T]$. Furthermore, by Lemma F.5, the quantity $\|\rho^\delta(t)\|_{\mathcal{L}^2(\Omega)}$ has a δ -free upper bound for $t \in [0, T_0]$. As a result, for almost all $t \in [0, T_0]$, the third integral in (F.59) tends to 0 as $\delta \rightarrow 0$. By iterating this argument T/T_0 times, we conclude that for almost all $t \in [0, T_0]$, the third integral in (F.59) tends to 0 as $\delta \rightarrow 0$, and the proof is complete. $\square$

At this point, we are ready to provide the proof of Theorem 5.3.

Proof of Theorem 5.3. By substituting z with $z^{1/2p}$ in Theorem 5.1, we have that with probability at least $1 - 1/z$

$$R_N(\mathbf{w}^k) \leq R^\delta(\rho_{k\varepsilon}^\delta) + z^{1/2p} \text{err}(N, d, \varepsilon, \delta) e^{C_* p \delta^{-(d+2)} T}, \quad (\text{F.67})$$

where $\text{err}(N, d, \varepsilon, \delta)$ is defined in (5.2). The risk $R^\delta(\rho_{k\varepsilon}^\delta)$ can be upper bounded as

$$\begin{aligned} R^\delta(\rho_{k\varepsilon}^\delta) &= \nu_0 \|f - K^\delta * \rho_{k\varepsilon}^\delta\|_{\mathcal{L}^2(\Omega)}^2 \\ &\leq \nu_0 \left(\|f - \rho_{k\varepsilon}^\delta\|_{\mathcal{L}^2(\Omega)} + \|K^\delta * \rho_{k\varepsilon}^\delta - \rho_{k\varepsilon}^\delta\|_{\mathcal{L}^2(\Omega)} \right)^2 \\ &= R(\rho_{k\varepsilon}^\delta) + \Delta_0(\delta, T, d), \end{aligned} \quad (\text{F.68})$$

where $\Delta_0(\delta, T, d) \rightarrow 0$ as $\delta \rightarrow 0$, since both $K^\delta * \rho_t^\delta$ and ρ_t^δ converge in $\mathcal{L}^2(\Omega)$ to ρ_t^∞ . Furthermore, by Theorem F.8,

$$\begin{aligned} R(\rho_{k\varepsilon}^\delta) &= 2F(\rho_{k\varepsilon}^\delta) + 2\tau S(\rho_{k\varepsilon}^\delta) \leq 2F(\rho_{k\varepsilon}^\delta) + 2\tau \log |\Omega| \\ &\leq 2F(\rho^*) + 2(F(\rho^\infty(0)) - F(\rho^*))e^{-2\alpha k\varepsilon} + 2\tau \log |\Omega| + \Delta(\delta, T, d) \\ &= 2F(\rho^\infty(0))e^{-2\alpha k\varepsilon} + 2(1 - e^{-2\alpha k\varepsilon})F(\rho^*) + 2\tau \log |\Omega| + \Delta(\delta, T, d), \end{aligned} \quad (\text{F.69})$$

where $\Delta(\delta, T, d) \rightarrow 0$ as $\delta \rightarrow 0$ and we recall that $|\Omega|$ denotes the volume of the set Ω .

Note that

$$F(\rho^*) \leq F(f) = -\tau S(f), \quad (\text{F.70})$$

since ρ^* is the minimizer of F . By combining (F.70) with (F.69), we deduce that

$$\begin{aligned} R(\rho_{k\varepsilon}^\delta) &\leq 2F(\rho^\infty(0))e^{-2\alpha k\varepsilon} + 2\tau \left(\log |\Omega| - (1 - e^{-2\alpha k\varepsilon})S(f) \right) + \Delta(\delta, T, d) \\ &= R(\rho^\infty(0))e^{-2\alpha k\varepsilon} + 2\tau \left(\log |\Omega| - (1 - e^{-2\alpha k\varepsilon})S(f) - S(\rho^\infty(0))e^{-2\alpha k\varepsilon} \right) \\ &\quad + \Delta(\delta, T, d) \\ &\leq R_N(\mathbf{w}^0)e^{-2\alpha k\varepsilon} + 2\tau \left(\log |\Omega| - (1 - e^{-2\alpha k\varepsilon})S(f) - S(\rho^\infty(0))e^{-2\alpha k\varepsilon} \right) \\ &\quad + z \text{err}(N, d, \varepsilon, \delta) e^{C_* p \delta^{-(d+2)} T} e^{-2\alpha k\varepsilon} + \Delta(\delta, T, d), \end{aligned} \quad (\text{F.71})$$

where in the last step we use again the result of Theorem 5.1 and the fact that $R(\rho^\infty(0)) - R^\delta(\rho^\infty(0))$ tends to 0 as $\delta \rightarrow 0$.

By optimizing over p in (F.67), we will set $\Delta_1(N, \varepsilon, T, d, z)$ as in (5.8). We also let $\Delta_2(\delta, T, d) = \Delta_0(\delta, T, d) + \Delta(\delta, T, d)$. Then, the result follows by combining (F.67), (F.68) and (F.71). $\square$

G Heat kernel in bounded domains with Neumann boundary

Given the domain $D \subseteq \mathbb{R}^d$ (compact, with $\mathcal{C}^2$ boundary ∂D), we denote by $G^D(\mathbf{x}, \mathbf{y}; t)$ the associated heat kernel, with Neumann boundary conditions. We collect here a few well known facts about this kernel (see, e.g., [Tay13, Section 6.1]).

The heat kernel can be defined as a function $G^D : D \times D \times \mathbb{R}_{>0}$ satisfying

$$\partial_t G^D(\mathbf{x}, \mathbf{y}; t) = \Delta_{\mathbf{y}} G^D(\mathbf{x}, \mathbf{y}; t), \quad (\text{G.1})$$

$$\langle \nabla_{\mathbf{y}} G^D(\mathbf{x}, \mathbf{y}; t), \mathbf{n}(\mathbf{y}) \rangle = 0 \quad \forall \mathbf{y} \in \partial D, \quad (\text{G.2})$$

$$G^D(\mathbf{x}, \cdot; t) \Rightarrow \delta_{\mathbf{x}}, \quad \text{as } t \rightarrow 0, \mathbf{x} \in D^\circ. \quad (\text{G.3})$$

We will also denote by $G(\mathbf{x}, \mathbf{y}; t)$ the heat kernel on $\mathbb{R}^d$, namely

$$G(\mathbf{x}, \mathbf{y}; t) \equiv \frac{1}{(4\pi t)^{d/2}} \exp \left\{ -\frac{\|\mathbf{x} - \mathbf{y}\|_2^2}{4t} \right\}. \quad (\text{G.4})$$

The probabilistic interpretation of G^D is as follows (see, e.g., [BGL13]). Let $\mathbb{E}_{\mathbf{x}}$ denote expectation with respect to a Brownian motion $\mathbf{X}_t$, with initial condition $\mathbf{X}_0 = \mathbf{x}$, and reflected at ∂D

(see Section C for definitions of this process, following [Tan79]). Then, for any bounded continuous function $\varphi : D \rightarrow \mathbb{R}$,

$$\mathbb{E}_{\mathbf{x}}\{\varphi(\mathbf{X}_t)\} = \int G^D(\mathbf{x}, \mathbf{y}; t) \varphi(\mathbf{y}) d\mathbf{y} \quad (\text{G.5})$$

$$\equiv G_\varphi^D(\mathbf{x}; t). \quad (\text{G.6})$$

Finally, G^D can be viewed as the kernel representation of the bounded operator $e^{t\Delta/2}$ in $\mathcal{L}^2(D, \text{Unif})$. We have

$$(e^{t\Delta/2}f)(\mathbf{y}) = \int f(\mathbf{x}) G^D(\mathbf{x}, \mathbf{y}; t) d\mathbf{x}. \quad (\text{G.7})$$

Hence $G^D(\mathbf{x}, \mathbf{y}; t)$ can be represented in terms of the eigenfunctions ϕ_k , and eigenvalues λ_k , of $-\Delta$,

$$G^D(\mathbf{x}, \mathbf{y}; t) = \sum_{k=0}^{\infty} e^{-\lambda_k t} \phi_k(\mathbf{x}) \phi_k(\mathbf{y}). \quad (\text{G.8})$$

Here $0 = \lambda_0 < \lambda_1 \leq \lambda_2 \leq \dots$, with $\lim_{k \rightarrow \infty} \lambda_k = \infty$, and $\phi_0(\mathbf{x}) = \mathbf{1}_D(\mathbf{x})/\text{Vol}(D)^{1/2}$.

Remark G.1. Since Δ is self-adjoint in $\mathcal{L}^2(D, \text{Unif})$, it follows that G^D is symmetric, namely $G^D(\mathbf{x}, \mathbf{y}, t) = G^D(\mathbf{y}, \mathbf{x}; t)$, and therefore it satisfies

$$\partial_t G^D(\mathbf{x}, \mathbf{y}; t) = \Delta_{\mathbf{x}} G^D(\mathbf{x}, \mathbf{y}; t), \quad (\text{G.9})$$

$$\langle \nabla_{\mathbf{x}} G^D(\mathbf{x}, \mathbf{y}; t), \mathbf{n}(\mathbf{x}) \rangle = 0 \quad \forall \mathbf{x} \in \partial D. \quad (\text{G.10})$$

Theorem G.1. *The Neumann heat kernel satisfies the following properties:*

1. *We have that*

$$G^D(\mathbf{x}, \mathbf{y}; t) = G(\mathbf{x}, \mathbf{y}; t) + G_R^D(\mathbf{x}, \mathbf{y}; t), \quad (\text{G.11})$$

where $G_R^D \in \mathcal{C}^\infty(D \times D \times \mathbb{R}_{\geq})$.

2. *For any $t > 0$, $G^D(\cdot, \cdot; t) \in \mathcal{C}^\infty(D \times D)$.*

3. *We have that, for a constant $C(D)$,*

$$\|\nabla G^D(\mathbf{x}, \cdot; t)\|_{\mathcal{L}^1(D)} \leq \frac{C(D)}{\sqrt{t}}. \quad (\text{G.12})$$

Proof. Substituting $G^D(\mathbf{x}, \mathbf{y}; t) = G(\mathbf{x}, \mathbf{y}; t) + G_R^D(\mathbf{x}, \mathbf{y}; t)$ into Eqs. (G.1) to (G.3) yields, for $\mathbf{x} \in D$,

$$\partial_t G_R^D(\mathbf{x}, \mathbf{y}; t) = \Delta_{\mathbf{y}} G_R^D(\mathbf{x}, \mathbf{y}; t), \quad (\text{G.13})$$

$$\langle \nabla_{\mathbf{y}} G_R^D(\mathbf{x}, \mathbf{y}; t), \mathbf{n}(\mathbf{y}) \rangle = -\langle \nabla_{\mathbf{y}} G(\mathbf{x}, \mathbf{y}; t), \mathbf{n}(\mathbf{y}) \rangle \quad \forall \mathbf{y} \in \partial D, \quad (\text{G.14})$$

$$G_R^D(\mathbf{x}, \mathbf{y}; 0) = 0, \quad \mathbf{x}, \mathbf{y} \in D^\circ. \quad (\text{G.15})$$

Thus G_R satisfies the heat equation in $D \times [0, T]$ and hence $(\mathbf{y}, t) \mapsto G_R^D(\mathbf{x}, \mathbf{y}; t)$ is $\mathcal{C}^\infty$ inside this domain (see, e.g., [Eva09, Chapter 2, Theorem 8], which refers to Dirichlet boundary condition, but applies equally well to the Neumann case). By symmetry, we have the claimed continuity in $(\mathbf{x}, \mathbf{y})$, thus proving point 1.

Claim 2 follows by the same decomposition.

Finally, claim 3 follows from Lemma 3.1 in [WY13]. $\square$

H Some useful technical lemmas

Lemma H.1 (Displacement convexity of quadratic functionals). *Let $U : \mathbb{R}^d \rightarrow \mathbb{R}^d$ be twice differentiable with $|U(\mathbf{x})| \leq C(1 + |\mathbf{x}|^2)$, $U(\mathbf{x}) = U(-\mathbf{x})$, and define $\mathcal{U} : \mathcal{P}_2(\mathbb{R}^d) \rightarrow \mathbb{R}$ by $\mathcal{U}(\rho) \equiv \int U(\mathbf{x} - \mathbf{x}') \rho(d\mathbf{x}) \rho(d\mathbf{x}')$. Then $\mathcal{U}$ is displacement convex if and only if U is convex.*

Proof. Proposition 7.4 in [San15] proves that convexity of U implies displacement convexity of $\mathcal{U}$. To prove the converse implication, let $\mathbf{x}, \boldsymbol{\delta} \in \mathbb{R}^d$, $\mathbf{x} \neq \mathbf{0}$ and consider the two probability distributions $\rho_0 = (\delta_{\mathbf{0}} + \delta_{\mathbf{x}})/2$ and $\rho_1 = (\delta_{\mathbf{0}} + \delta_{\mathbf{x}+\boldsymbol{\delta}})/2$. For $|\boldsymbol{\delta}| < |\mathbf{x}|$, the geodesic path connecting these distribution is $\rho_t = (\delta_{\mathbf{0}} + \delta_{\mathbf{x}+t\boldsymbol{\delta}})/2$, $t \in [0, 1]$. Substituting in the definition of $\mathcal{U}$, we get

$$\mathcal{U}(\rho_t) = \frac{1}{2} U(\mathbf{0}) + \frac{1}{2} U(\mathbf{x} + t\boldsymbol{\delta}) \quad (\text{H.1})$$

$$= \mathcal{U}(\rho_0) + \frac{t}{2} \langle \nabla U(\mathbf{x}), \boldsymbol{\delta} \rangle + \frac{t^2}{4} \langle \boldsymbol{\delta}, \nabla^2 U(\mathbf{x}) \boldsymbol{\delta} \rangle + o(t^2). \quad (\text{H.2})$$

Hence, displacement convexity implies $\langle \boldsymbol{\delta}, \nabla^2 U(\mathbf{x}) \boldsymbol{\delta} \rangle \geq 0$. Since this holds for all $|\boldsymbol{\delta}| < |\mathbf{x}|$, we obtain $\nabla^2 U(\mathbf{x}) \succeq \mathbf{0}$ for all $\mathbf{x} \neq \mathbf{0}$, which in turns imply that U is convex (by a continuity argument, it is sufficient to lower bound the Hessian everywhere except at a point). $\square$

Lemma H.2 (A Gronwall type inequality [Bih56]). *Let $u : [0, T] \rightarrow \mathbb{R}_+$ be a continuous function that satisfies the inequality*

$$u(t) \leq A + \int_0^t \Psi(s) \omega(u(s)) ds, \quad t \in [0, T],$$

where $A \geq 0$, $\Psi : [0, T] \rightarrow \mathbb{R}_+$ is continuous and $\omega : \mathbb{R}_+ \rightarrow \mathbb{R}_+$ is continuous and monotone-increasing. Then, the following holds

$$u(t) \leq \Phi^{-1} \left(\Phi(A) + \int_0^t \Psi(s) ds \right), \quad t \in [0, T],$$

with $\Phi : \mathbb{R} \mapsto \mathbb{R}$ given by

$$\Phi(u) \equiv \int_{u_0}^u \frac{ds}{\omega(s)}, \quad u \in \mathbb{R}, \quad u_0 \equiv \omega(A).$$

Remark H.1. To derive Equation (F.45), we use Lemma H.2 with $\omega(u) = u^2$, $\Psi(s) = \bar{C}_m C_3$, $A = \bar{C}_m C_4 + \bar{C}_m \tau M^2 C_a s t T_n$.

Lemma H.3 (Gagliardo-Nirenberg interpolation inequality, cf. Theorem 1.5.2 of [CM12]). *Fix $1 \leq q, r \leq \infty$ and m a positive integer. Let $u \in \mathcal{L}^q(\Omega) \cap \mathcal{L}^r(\Omega)$ and $\nabla^{\otimes m} u \in \mathcal{L}^p(\Omega)$. For integer j , $0 \leq j \leq m$, and $\theta \in [j/m, 1]$ (with the exception $\theta \neq 1$ if $m - j - d/2$ is a non-negative integer), define p by*

$$\frac{1}{p} = \frac{j}{d} + \theta \left(\frac{1}{r} - \frac{m}{d} \right) + \frac{1 - \theta}{q}.$$

Then $\nabla^{\otimes j} u \in \mathcal{L}^p(\Omega)$ and satisfies

$$\|\nabla^{\otimes j} u\|_p \leq C \|\nabla^{\otimes m} u\|_r^\theta \|u\|_q^{1-\theta} + C_1 \|u\|_s.$$

with finite arbitrary $1 \leq s \leq \max(r, q)$ and $C > 0$ and $C_1 \geq 0$ are independent of u . The constant C is independent of Ω , while $C_1 \rightarrow 0$ as $|\Omega| \rightarrow \infty$. In particular, the choice $C_1 = 0$ is admissible if $\Omega = \mathbb{R}^d$.

References

- [AB09] Martin Anthony and Peter L. Bartlett, *Neural network learning: Theoretical foundations*, Cambridge University Press, 2009. 2, 6
- [AGS08] Luigi Ambrosio, Nicola Gigli, and Giuseppe Savaré, *Gradient flows: in metric spaces and in the space of probability measures*, Springer Science & Business Media, 2008. 10
- [Bac17] Francis Bach, *Breaking the curse of dimensionality with convex neural networks*, The Journal of Machine Learning Research **18** (2017), no. 1, 629–681. 3
- [Bar93] Andrew R. Barron, *Universal approximation bounds for superpositions of a sigmoidal function*, IEEE Transactions on Information theory **39** (1993), no. 3, 930–945. 6
- [Bar98] Peter L. Bartlett, *The sample complexity of pattern classification with neural networks: the size of the weights is more important than the size of the network*, IEEE Transactions on Information Theory **44** (1998), no. 2, 525–536. 6, 24
- [BGL13] Dominique Bakry, Ivan Gentil, and Michel Ledoux, *Analysis and geometry of markov diffusion operators*, vol. 348, Springer Science & Business Media, 2013. 64
- [Bih56] Imre Bihari, *A generalization of a lemma of Bellman and its application to uniqueness problems of differential equations*, Acta Mathematica Hungarica **7** (1956), no. 1, 81–94. 66
- [BJW18] Ainesh Bakshi, Rajesh Jayaram, and David P Woodruff, *Learning two layer rectified neural networks in polynomial time*, arXiv:1811.01885 (2018). 5
- [BRV⁺06] Yoshua Bengio, Nicolas L. Roux, Pascal Vincent, Olivier Delalleau, and Patrice Marcotte, *Convex neural networks*, Advances in Neural Information Processing Systems, 2006, pp. 123–130. 6
- [BY03] Peter Bühlmann and Bin Yu, *Boosting with the L2 loss: regression and classification*, Journal of the American Statistical Association **98** (2003), no. 462, 324–339. 2
- [CB18] Lenaïc Chizat and Francis Bach, *On the global convergence of gradient descent for over-parameterized models using optimal transport*, Advances in Neural Information Processing Systems, 2018. 4, 5, 6
- [CFG14] Emmanuel J. Candès and Carlos Fernandez-Granda, *Towards a mathematical theory of super-resolution*, Communications on Pure and Applied Mathematics **67** (2014), no. 6, 906–956. 2
- [CJM⁺01] José A. Carrillo, Ansgar Jüngel, Peter A. Markowich, Giuseppe Toscani, and Andreas Unterreiter, *Entropy dissipation methods for degenerate parabolic problems and generalized sobolev inequalities*, Monatshefte für Mathematik **133** (2001), no. 1, 1–82. 4, 5, 23, 61
- [CM12] Pascal Cherrier and Albert Milani, *Linear and quasi-linear evolution equations in Hilbert spaces*, Graduate studies in mathematics, American Mathematical Soc., 2012. 66

- [CMV03] José A. Carrillo, Robert J. McCann, and Cédric Villani, *Kinetic equilibration rates for granular media and related equations: entropy dissipation and mass transportation estimates*, *Revista Matemática Iberoamericana* **19** (2003), no. 3, 971–1018. [4](#), [5](#), [23](#)
- [CMV06] ———, *Contractions in the 2-Wasserstein length space and thermalization of granular media*, *Archive for Rational Mechanics and Analysis* **179** (2006), no. 2, 217–263. [4](#), [5](#), [23](#)
- [CS16] Yining Chen and Richard J Samworth, *Generalized additive and index models with shape constraints*, *Journal of the Royal Statistical Society: Series B (Statistical Methodology)* **78** (2016), no. 4, 729–754. [11](#)
- [CST00] Nello Cristianini and John Shawe-Taylor, *An introduction to support vector machines and other kernel-based learning methods*, Cambridge University Press, 2000. [2](#)
- [CT06] Thomas M Cover and Joy A Thomas, *Elements of information theory*, John Wiley & Sons, 2006. [63](#)
- [Cyb89] George Cybenko, *Approximation by superpositions of a sigmoidal function*, *Mathematics of control, signals and systems* **2** (1989), no. 4, 303–314. [6](#)
- [Dob79] Roland L’vovich Dobrushin, *Vlasov equations*, *Functional Analysis and Its Applications* **13** (1979), no. 2, 115–123. [6](#)
- [Don92] David L. Donoho, *Superresolution via sparsity constraints*, *SIAM Journal on Mathematical Analysis* **23** (1992), no. 5, 1309–1331. [2](#)
- [DZPS18] Simon S Du, Xiyu Zhai, Barnabas Poczos, and Aarti Singh, *Gradient descent provably optimizes over-parameterized neural networks*, *arXiv:1810.02054* (2018). [5](#)
- [Eva09] Lawrence C. Evans, *Partial differential equations*, Springer, 2009. [65](#)
- [FP08] Alessio Figalli and Robert Philipowski, *Convergence to the viscous porous medium equation and propagation of chaos*, *ALEA Lat. Am. J. Probab. Math. Stat* **4** (2008), 185–203. [22](#)
- [Fri01] Jerome H. Friedman, *Greedy function approximation: a gradient boosting machine*, *Annals of Statistics* (2001), 1189–1232. [2](#)
- [HD13] Lauren A Hannah and David B Dunson, *Multivariate convex regression with adaptive partitioning*, *The Journal of Machine Learning Research* **14** (2013), no. 1, 3261–3294. [11](#)
- [LS84] Pierre-Louis Lions and Alain-Sol Sznitman, *Stochastic differential equations with reflecting boundary conditions*, *Communications on Pure and Applied Mathematics* **37** (1984), no. 4, 511–537. [21](#)
- [LSU88] Olga Aleksandrovna Ladyzhenskaia, Vsevolod Alekseevich Solonnikov, and Nina N. Ural’tseva, *Linear and quasi-linear equations of parabolic type*, vol. 23, American Mathematical Society, 1988. [28](#), [45](#)

- [LY17] Yuanzhi Li and Yang Yuan, *Convergence analysis of two-layer neural networks with relu activation*, Advances in Neural Information Processing Systems, 2017, pp. 597–607. 5
- [MBM⁺18] Song Mei, Yu Bai, Andrea Montanari, et al., *The landscape of empirical risk for non-convex losses*, The Annals of Statistics **46** (2018), no. 6A, 2747–2774. 24
- [McC97] Robert J McCann, *A convexity principle for interacting gases*, Advances in mathematics **128** (1997), no. 1, 153–179. 10
- [MMM19] Song Mei, Theodor Misiakiewicz, and Andrea Montanari, *Mean-field theory of two-layers neural networks: dimension-free bounds and kernel limit*, Conference on Learning Theory (COLT), 2019. 6, 24
- [MMN18] Song Mei, Andrea Montanari, and Phan-Minh Nguyen, *A mean field view of the landscape of two-layer neural networks*, Proceedings of the National Academy of Sciences (2018). 4, 5, 6, 19, 21, 24, 31, 38, 43, 44, 45
- [NS17] Atsushi Nitanda and Taiji Suzuki, *Stochastic particle gradient descent for infinite ensembles*, arXiv:1712.05438 (2017). 6
- [Oel01] Karl Oelschläger, *A sequence of integro-differential equations approximating a viscous porous medium equation*, Zeitschrift für Analysis und ihre Anwendungen **20** (2001), no. 1, 55–91. 55, 56, 57
- [Oel02] ———, *Simulation of the solution of a viscous porous medium equation by a particle method*, SIAM Journal on Numerical Analysis **40** (2002), no. 5, 1716–1762. 22
- [Phi07] Robert Philipowski, *Interacting diffusions approximating the porous medium equation and propagation of chaos*, Stochastic Processes and their Applications **117** (2007), no. 4, 526–538. 22
- [PS91] Jooyoung Park and Irwin W. Sandberg, *Universal approximation using radial-basis-function networks*, Neural computation **3** (1991), no. 2, 246–257. 3
- [Ros62] Frank Rosenblatt, *Principles of neurodynamics*, Spartan Book, 1962. 2
- [RR08] Ali Rahimi and Benjamin Recht, *Random features for large-scale kernel machines*, Advances in neural information processing systems, 2008, pp. 1177–1184. 2
- [RVE18] Grant M. Rotskoff and Eric Vanden-Eijnden, *Neural networks as interacting particle systems: Asymptotic convexity of the loss landscape and universal scaling of the approximation error*, Advances in Neural Information Processing Systems, 2018. 4, 6
- [RW94] L. Chris G. Rogers and David Williams, *Diffusions, Markov processes and martingales: Volume 2, Itô calculus*, vol. 2, Cambridge university press, 1994. 36
- [San15] Filippo Santambrogio, *Optimal transport for applied mathematicians: Calculus of variations, PDEs, and modeling*, vol. 87, Birkhäuser, 2015. 10, 66
- [Sch03] Robert E. Schapire, *The boosting approach to machine learning: An overview*, Nonlinear estimation and classification, Springer, 2003, pp. 149–171. 2

- [SJL18] Mahdi Soltanolkotabi, Adel Javanmard, and Jason D. Lee, *Theoretical insights into the optimization landscape of over-parameterized shallow neural networks*, IEEE Transactions on Information Theory (2018). 5
- [Slo94] Slomiński, Leszek, *On approximation of solutions of multidimensional SDE's with reflecting boundary conditions*, Stochastic processes and their Applications **50** (1994), no. 2, 197–219. 38, 40
- [Slo01] ———, *Euler's approximations of solutions of SDEs with reflecting boundary*, Stochastic processes and their applications **94** (2001), no. 2, 317–337. 38, 48
- [SS18] Justin Sirignano and Konstantinos Spiliopoulos, *Mean field analysis of neural networks*, arXiv:1805.01053 (2018). 4, 6
- [Szn91] Alain-Sol Sznitman, *Topics in propagation of chaos*, Ecole d'été de probabilités de Saint-Flour XIX—1989, Springer, 1991, pp. 165–251. 6, 38
- [Tan79] Hiroshi Tanaka, *Stochastic differential equations with reflecting boundary condition in convex regions*, Hiroshima Mathematical Journal **9** (1979), no. 1, 163–177. 21, 35, 37, 65
- [Tay13] Michael Taylor, *Partial differential equations I: Basic theory*, vol. 115, Springer Science & Business Media, 2013. 64
- [Tho13] James William Thomas, *Numerical partial differential equations: finite difference methods*, vol. 22, Springer Science & Business Media, 2013. 12
- [Tia17] Yuandong Tian, *Symmetry-breaking convergence analysis of certain two-layered neural networks with ReLU nonlinearity*, Workshop at International Conference on Learning Representation (ICLR), 2017. 5
- [Váz07] Juan Luis Vázquez, *The porous medium equation: mathematical theory*, Oxford University Press, 2007. 10, 28
- [Vil08] Cédric Villani, *Optimal transport: old and new*, vol. 338, Springer Science & Business Media, 2008. 10
- [WLLM18] Colin Wei, Jason D Lee, Qiang Liu, and Tengyu Ma, *On the margin theory of feedforward neural networks*, arXiv:1810.05369 (2018). 5
- [WY13] Feng-Yu Wang and Lixin Yan, *Gradient estimate on convex domains and applications*, Proceedings of the American Mathematical Society **141** (2013), no. 3, 1067–1081. 65
- [ZSJ⁺17] Kai Zhong, Zhao Song, Prateek Jain, Peter L. Bartlett, and Inderjit S. Dhillon, *Recovery guarantees for one-hidden-layer neural networks*, International Conference on Machine Learning, 2017, pp. 4140–4149. 5