

A Dynamical Central Limit Theorem for Shallow Neural Networks

Zhengdao Chen^a, Grant M. Rotskoff^{a, b}, Joan Bruna^{a, c}, and Eric Vanden-Eijnden^a

^aCourant Institute of Mathematical Sciences, New York University, New York

^bDepartment of Chemistry, Stanford University, California

^cCenter for Data Science, New York University, New York

November 20, 2020

Abstract

Recent theoretical works have characterized the dynamics of wide shallow neural networks trained via gradient descent in an asymptotic mean-field limit when the width tends towards infinity. At initialization, the random sampling of the parameters leads to deviations from the mean-field limit dictated by the classical Central Limit Theorem (CLT). However, since gradient descent induces correlations among the parameters, it is of interest to analyze how these fluctuations evolve. Here, we use a dynamical CLT to prove that the asymptotic fluctuations around the mean limit remain bounded in mean square throughout training. The upper bound is given by a Monte-Carlo resampling error, with a variance that depends on the 2-norm of the underlying measure, which also controls the generalization error. This motivates the use of this 2-norm as a regularization term during training. Furthermore, if the mean-field dynamics converges to a measure that interpolates the training data, we prove that the asymptotic deviation eventually vanishes in the CLT scaling. We also complement these results with numerical experiments.

1 Introduction

Theoretical analyses of neural networks aim to understand their computational and statistical advantages seen in practice. On the computation side, the training of neural networks often succeed despite being a non-convex optimization problem known to be hard in certain settings [20, 31, 41]. On the statistics side, neural networks often generalize well despite having large numbers of parameters [8, 70]. In this context, the notion of *over-parametrization* has been useful, by providing insights into the optimization and generalization properties as the network widths tend to infinity [2, 4, 21, 36, 38, 63, 67]. In particular, under appropriate scaling, one can view shallow (a.k.a. single-hidden-layer or two-layer) networks as interacting particle systems that admit a mean-field limit. Their training dynamics can then be studied as Wasserstein Gradient Flows [12, 47, 51, 58], leading to global convergence guarantees in the mean-field limit under certain assumptions. On the statistics

Correspondence to: zc1216@nyu.edu, rotskoff@stanford.edu, bruna@cims.nyu.edu and eve2@cims.nyu.edu.

side, such an approach lead to powerful generalization guarantees for learning high-dimensional functions with hidden low-dimensional structures, as compared to learning in Reproducing Kernel Hilbert Spaces (RKHS) [5, 30]. However, since ultimately we are concerned with neural networks of finite width, it is key to study the deviation of finite-width networks from their infinite-width limits, and how it scales with the width m . At the random initial state, neurons do not interact and therefore a standard Monte-Carlo (MC) argument shows that the fluctuations in the underlying measure scale as $m^{-1/2}$, which we refer to as the Central Limit Theorem (CLT) scaling. As optimization introduces complex dependencies among the parameters, the key question is to understand how the fluctuation evolves during training. To make this investigation tractable, we aim to obtain insight on an asymptotic scale as the width grows, and focus on the evolution in time. An application of Grönwall’s inequality shows that this asymptotic deviation remains bounded at all finite time [46], but the dependence on time is exponential, making it difficult to assess the long-time behavior.

The main focus of this paper is to investigate this question in-depth, by analyzing the interplay between the deviations from the mean-field limit and the gradient flow dynamics. First, we prove a dynamical CLT to capture how the fluctuations away from the mean-field limit evolve as a function of training time to show that the fluctuations remain on the initial $m^{-1/2}$ -scale for all finite times. Next, we examine the long-time behavior of the fluctuations, proving that, in several scenarios, the long-time fluctuations are controlled by the error of Monte-Carlo resampling from the limiting measure. We focus on two main setups relevant for supervised learning and scientific computing: the unregularized case with global convergence of mean-field gradient flows to minimizers that interpolate the data, and the regularized case where the limiting measure has atomic support and is nondegenerate. In the former setup, we prove particularly that the fluctuations eventually vanish in the CLT scaling. These asymptotic predictions are complemented by empirical results in a teacher-student model.

Related Works: This paper continues the line of work initiated in [12, 47, 51, 58] that studies optimization of over-parameterized shallow neural networks under the mean-field scaling. Global convergence for the unregularized setting is discussed in [46, 47, 51, 58]. In the regularized setting, [12] establishes global convergence in the mean-field limit under specific homogeneity conditions on the neuron activation. Other works that study asymptotic properties of wide neural networks include [1, 6, 23, 28, 29, 34, 35, 43, 69], notably investigating the transition between the so-called *lazy* and *active* regimes [14], corresponding respectively to linear versus nonlinear learning. Our focus is on the dynamics under the mean-field scaling, which encompasses the active, nonlinear regime.

A relevant work concerning the sparse optimization of measures is [11], where under a different metric for gradient flow and additional assumptions on the nature of the minimizer, it can be established that fluctuations vanish for sufficiently large m . Our results are only asymptotic in m but apply to broader settings in the context of shallow neural networks. Concerning the next-order deviations of finite neural networks from their mean-field limit, [51] show that the scale of fluctuations is below that of MC resampling for unregularized problems using non-rigorous arguments. [60] provides a CLT for the fluctuations at finite time under stochastic gradient descent (SGD) and proves that the fluctuations decay in time in the case where there is a single critical point in the parameter space. Our focus is on the long-time behavior of the fluctuations in more general settings. Another relevant topic is the propagation of chaos in McKean-Vlasov systems, which study the deviations of randomly-forced interacting particle systems from their infinite-particle limits [7, 10, 65, 66]. In particular, a line of work provides uniform-in-time bounds to the fluctuations in various settings [16, 19, 22, 55, 56], but the conditions are not applicable to shallow neural networks. Concurrently to our work, [17] studies quantitative propagation of chaos of shallow

neural networks trained by SGD, but the bound grows exponentially in time, and therefore cannot address the long-time behavior of the fluctuations.

Learning with neural networks exhibits the phenomenon that generalization error can decrease with the level of overparameterization [8, 64]. [48] proposes a bias-variance decomposition that contains a variance term initialization in optimization. They show in experiments that this term decreases as the width of the network increases, and justifies this theoretically under the strong assumption that model parameters remain Gaussian-distributed in the components that are irrelevant for the task, which does not hold in the scenario we consider, for example. [27] provides scaling arguments for the dependence of this term on the width of the network. Our work provides a more rigorous analysis of the dependence of this term on the width of the network and training time.

2 Background

2.1 Shallow Neural Networks and the Integral Representation

On a data space $\Omega \subseteq \mathbb{R}^d$, we consider parameterized models of the following form

$$f^{(m)}(\mathbf{x}) = \frac{1}{m} \sum_{i=1}^m \varphi(\boldsymbol{\theta}_i, \mathbf{x}), \quad (1)$$

where $\mathbf{x} \in \Omega$, $\{\boldsymbol{\theta}_i\}_{i=1}^m \subseteq D$ is the set of model parameters, and $\varphi : D \times \Omega \rightarrow \mathbb{R}$ is the activation function. Of particular interest are shallow neural network models, which admit a more specific form:

Assumption 2.1 (Shallow neural networks setting). $D = \mathbb{R} \times \hat{D}$, $\boldsymbol{\theta} = (c, \mathbf{z}) \in D$, and $\varphi(\boldsymbol{\theta}, \mathbf{x}) = c\hat{\varphi}(\mathbf{z}, \mathbf{x})$ with $\hat{\varphi} : \hat{D} \times \Omega \rightarrow \mathbb{R}$. Thus, (1) can be rewritten as $f^{(m)}(\mathbf{x}) = \frac{1}{m} \sum_{i=1}^m c_i \hat{\varphi}(\mathbf{z}_i, \mathbf{x})$.

As many of our results hold for general models of the form (1), we will invoke Assumption 2.1 only when needed. We shall also assume the following:

Assumption 2.2. Ω is compact; D is a Euclidean space (or a subset thereof); $\varphi(\boldsymbol{\theta}, \mathbf{x})$ is twice differentiable in $\boldsymbol{\theta}$; $\nabla_{\boldsymbol{\theta}} \nabla_{\boldsymbol{\theta}} \varphi(\boldsymbol{\theta}, \mathbf{x})$ is Lipschitz in $\boldsymbol{\theta}$, uniformly in $\mathbf{x}$.

The regularity assumptions are standard in the literature [10, 11, 37]. We note that they are not satisfied by ReLU units (i.e., $\hat{\varphi}(\mathbf{z}, \mathbf{x}) = \max\{0, \langle \mathbf{a}, \mathbf{x} \rangle + b\}$, where $\mathbf{z} = (\mathbf{a}, b)^\top$, with $\mathbf{a} \in \mathbb{R}^d$ and $b \in \mathbb{R}$), though prior work [12, 13] has considered differentiable approximations of these models.

As observed in [12, 24, 47, 51, 58], a model of the form (1) can be expressed in integral form in terms of a probability measure over D as $f^{(m)} = f[\mu^{(m)}]$, where we define

$$f[\mu](\mathbf{x}) = \int_D \varphi(\boldsymbol{\theta}, \mathbf{x}) \mu(d\boldsymbol{\theta}), \quad (2)$$

and $\mu^{(m)}$ is the empirical measure of the parameters $\{\boldsymbol{\theta}_i\}_{i=1}^m$:

$$\mu^{(m)}(d\boldsymbol{\theta}) = \frac{1}{m} \sum_{i=1}^m \delta_{\boldsymbol{\theta}_i}(d\boldsymbol{\theta}). \quad (3)$$

Suppose we are given a dataset $\{(\mathbf{x}_l, y_l)\}_{l=1}^n$, which can be represented by an empirical data measure $\hat{\nu} = \frac{1}{n} \sum_{l=1}^n \delta_{\mathbf{x}_l}$, and $y_l = f_*(\mathbf{x}_l)$ are generated by an target function f_* that we wish to

estimate using least-squares regression. A canonical approach to this regression task is to consider an Empirical Risk Minimization (ERM) problem of the form

$$\min_{\mu \in \mathcal{P}(D)} \mathcal{L}(\mu) \quad \text{with} \quad \mathcal{L}(\mu) := \frac{1}{2} \|f[\mu] - f_*\|_{\hat{\nu}}^2 + \lambda \int_D r(\boldsymbol{\theta}) \mu(d\boldsymbol{\theta}) . \quad (4)$$

where $\mathcal{P}(D)$ is the space of probability measures on D , $\|f - f_*\|_{\hat{\nu}}^2 := \int_{\Omega} |f(\mathbf{x}) - f_*(\mathbf{x})|^2 \hat{\nu}(d\mathbf{x})$ denotes the L_2 function reconstruction error averaged over the data, and $\lambda \int_D r(\boldsymbol{\theta}) \mu(d\boldsymbol{\theta})$ is some optional regularization term. While we can allow r to be a general convex function, in Section 3.3 we will motivate a choice of r in the shallow neural networks setting that is related to the variation norm [5] or Barron norm [44] of functions.

2.2 Approximation and Optimization with a Finite Number of Neurons

Integral representations with a probability measure such as those defined in (2) are amenable to efficient approximation in high dimensions via Monte-Carlo sampling. Namely, if the parameters $\boldsymbol{\theta}_i$ in $f^{(m)}$ are drawn i.i.d. from an underlying measure μ on D , then by the Law of Large Numbers (LLN), the resulting empirical measure $\mu^{(m)}$ converges μ almost surely, and moreover,

$$\mathbb{E}_{\mu^{(m)}} \|f[\mu^{(m)}] - f[\mu]\|_{\hat{\nu}}^2 = \frac{1}{m} \left(\int_D \|\varphi(\boldsymbol{\theta}, \cdot)\|_{\hat{\nu}}^2 \mu(d\boldsymbol{\theta}) - \|f[\mu]\|_{\hat{\nu}}^2 \right) , \quad (5)$$

Such a Monte-Carlo estimator showcases the benefit of normalized integral representations for high-dimensional approximation, as the ambient dimension appears in the rate of approximation only through the term $\int_D \|\varphi(\boldsymbol{\theta}, \cdot)\|_{\hat{\nu}}^2 \mu(d\boldsymbol{\theta})$. In the case of shallow neural networks, this is connected to the variation norm or Barron norm of the function we wish to approximate [5, 44] (see Section 3.3 for details).

While the Monte-Carlo sampling strategy above can be seen as a ‘static’ approximation of a function representable as (2), it also gives rise to an efficient algorithm to optimize (4). Indeed, in terms of the empirical distribution $\mu^{(m)}$, the loss $\mathcal{L}(\mu^{(m)})$ becomes a function of the parameters $\{\boldsymbol{\theta}_i\}_{i=1}^m$, which we can seek to minimize by adjusting the parameters:

$$L(\boldsymbol{\theta}_1, \dots, \boldsymbol{\theta}_m) = \frac{1}{2} \|f^{(m)} - f_*\|_{\hat{\nu}}^2 + \frac{\lambda}{m} \sum_{i=1}^m r(\boldsymbol{\theta}_i) . \quad (6)$$

In the shallow neural network setting, with suitable choices of the function r , the regularization term corresponds to *weight decay* over the parameters.

2.3 From Particle to Wasserstein Gradient Flows

Expanding (6), we get

$$L(\boldsymbol{\theta}_1, \dots, \boldsymbol{\theta}_m) = C_{f_*} - \frac{1}{m} \sum_{i=1}^m F(\boldsymbol{\theta}_i) + \frac{1}{2m^2} \sum_{i,j=1}^m K(\boldsymbol{\theta}_i, \boldsymbol{\theta}_j), \quad (7)$$

where we have defined $C_f = \frac{1}{2} \|f\|_{\hat{\nu}}^2$, and

$$F(\boldsymbol{\theta}) = \int_{\Omega} f_*(\mathbf{x}) \varphi(\boldsymbol{\theta}, \mathbf{x}) \hat{\nu}(d\mathbf{x}) - \lambda r(\boldsymbol{\theta}), \quad K(\boldsymbol{\theta}, \boldsymbol{\theta}') = \int_{\Omega} \varphi(\boldsymbol{\theta}, \mathbf{x}) \varphi(\boldsymbol{\theta}', \mathbf{x}) \hat{\nu}(d\mathbf{x}) . \quad (8)$$

Performing GD on L amounts to discretizing in time the following ODE system that governs the evolution of for $\{\theta_i\}_{i=1}^m$:

$$\dot{\theta}_i = -m \partial_{\theta_i} L(\theta_1 \dots \theta_m) = \nabla F(\theta_i) - \frac{1}{m} \sum_{j=1}^m \nabla K(\theta_i, \theta_j) =: -\nabla V(\theta_i, \mu_t^{(m)}). \quad (9)$$

where we defined the potential

$$V(\theta, \mu) = -F(\theta) + \int_D K(\theta, \theta') \mu(d\theta'). \quad (10)$$

Heuristically, the ‘particles’ θ_i perform GD according to the potential $V(\theta, \mu_t^{(m)})$ which itself evolves, depending on the particles positions through their empirical measure. Such dynamics can also be expressed in terms of the empirical measure via the *continuity equation*:

$$\partial_t \mu_t^{(m)} = \nabla \cdot (\nabla V(\theta, \mu_t^{(m)}) \mu_t^{(m)}) \quad (11)$$

This equation should be understood in the weak sense by testing it against continuous functions $\chi : D \rightarrow \mathbb{R}$, and it can be interpreted as the gradient flow on the loss defined in (4) under the 2-Wasserstein metric [12, 47, 51, 58]. This insight provides powerful analytical tools to understand convergence properties, by considering the mean-field limit when $m \rightarrow \infty$.

2.4 Law of Large Numbers and Mean-Field Gradient Flow

From now on, we assume that the particle gradient flow is initialized in the following way:

Assumption 2.3. *The ODE (9) is solved for the initial condition $\theta_i(0) = \theta_i^0$, with θ_i^0 drawn i.i.d. from a compactly supported measure $\mu_0 \in \mathcal{P}(D)$ for each $i = 1, \dots, m$. Hence, $\mu_0^{(m)}(d\theta) = \frac{1}{m} \sum_{i=1}^m \delta_{\theta_i^0}(d\theta)$.*

We use $\mathbb{P}_0$ to denote the probability measure associated with the set $\{\theta_i^0\}_{i \in \mathbb{N}}$ with each θ_i^0 drawn i.i.d. from μ_0 , and use $\mathbb{E}_0$ to denote the expectation under $\mathbb{P}_0$. The Law of Large Numbers (LLN) indicates that $\mathbb{P}_0$ -almost surely, $\mu_t^{(m)} \rightarrow \mu_t$ as $m \rightarrow \infty$, where μ_t satisfies the mean-field gradient flow [12, 47, 52, 61]:

$$\partial_t \mu_t = \nabla \cdot (\nabla V(\theta, \mu_t) \mu_t), \quad \mu_{t=0} = \mu_0. \quad (12)$$

The solution to this equation can be understood via the representation formula

$$\int_D \chi(\theta) \mu_t(d\theta) = \int_D \chi(\Theta_t(\theta)) \mu_0(d\theta), \quad (13)$$

where χ is a continuous test function $\chi : D \rightarrow \mathbb{R}$ and $\Theta_t : D \rightarrow D$ is the *characteristic flow* associated with (11), which in direct analogy with (9) solves

$$\dot{\Theta}_t(\theta) = -\nabla V(\Theta_t(\theta), \mu_t), \quad \Theta_0(\theta) = \theta. \quad (14)$$

Using expression (10) for V as well as (13), this equation can be written in closed form explicitly as

$$\dot{\Theta}_t(\theta) = \nabla F(\Theta_t(\theta)) - \int_D \nabla K(\Theta_t(\theta), \Theta_t(\theta')) \mu_0(d\theta'), \quad \Theta_0(\theta) = \theta. \quad (15)$$

It is easy to see that this equation is itself a gradient flow since it is the continuous-time limit of a proximal scheme (mirror descent), which we state as:

Proposition 2.4. Given $\bar{\Theta}_0(\theta) = \theta$ and $\tau > 0$, for $p \in \mathbb{N}$ let $\Theta_{p\tau}$ be specified via

$$\bar{\Theta}_{p\tau} \in \operatorname{argmin} \left(\frac{1}{2\tau} \|\Theta - \bar{\Theta}_{(p-1)\tau}\|_0^2 + \mathcal{E}(\Theta) \right), \quad (16)$$

where we defined

$$\|\Theta\|_0^2 = \int_D |\Theta(\theta)|^2 \mu_0(d\theta) \quad (17)$$

and

$$\mathcal{E}(\Theta) = - \int_D F(\Theta(\theta)) \mu_0(d\theta) + \frac{1}{2} \int_D K(\Theta(\theta), \Theta(\theta')) \mu_0(d\theta) \mu_0(d\theta'). \quad (18)$$

Then

$$\lim_{\tau \rightarrow 0} \bar{\Theta}_{\lfloor t/\tau \rfloor \tau} = \Theta_t \quad \mu_0\text{-almost surely}, \quad (19)$$

where Θ_t solves (15).

2.5 Long-Time Properties of the Mean-Field Gradient Flow

In the shallow neural networks setting, a series of earlier works [12, 47, 51, 58] has established that under certain assumptions μ_t will converge to a global minimizer of the loss functional $\mathcal{L}$. In particular, [12] studies global convergence for the regularized loss $\mathcal{L}$ under homogeneity assumptions on $\hat{\varphi}$, and [50] considers modified dynamics using *double-lifting*. Here, to study the long time behavior of the fluctuations, we will often work with the following weaker assumptions:

Assumption 2.5. The solution to (15) exists for all time, and has a limit:

$$\Theta_t \rightarrow \Theta_\infty \quad \mu_0\text{-almost surely as } t \rightarrow \infty. \quad (20)$$

Assumption 2.6. The limiting Θ_∞ is a local minimizer of (18).

With these assumptions, we have

Proposition 2.7. Under Assumptions 2.3 and 2.5, we have

$$\cup_{t \geq 0} \operatorname{supp} \mu_t = \cup_{t \geq 0} \{\Theta_t(\theta) : \theta \in \operatorname{supp} \mu_0\} \text{ is compact}, \quad (21)$$

and $\mu_t \rightharpoonup \mu_\infty$ weakly as $t \rightarrow \infty$, with μ_∞ satisfying

$$\int_D \chi(\theta) \mu_\infty(d\theta) = \int_D \chi(\Theta_\infty(\theta)) \mu_0(d\theta), \quad (22)$$

for all continuous test function $\chi : D \rightarrow \mathbb{R}$. Additionally, if Assumption 2.6 also holds, then

$$\nabla \nabla V(\Theta_\infty(\theta), \mu_\infty) \text{ is positive semidefinite for } \mu_0\text{-almost all } \theta \quad (23)$$

We prove this proposition in Appendix B. Here, $\nabla \nabla V(\Theta_\infty(\theta), \mu_\infty)$ denotes

$$\nabla \nabla V(\Theta_\infty(\theta), \mu_\infty) = -\nabla \nabla F(\Theta_\infty(\theta)) + \int_D \nabla \nabla K(\Theta_\infty(\theta), \Theta_\infty(\theta')) \mu_0(d\theta'), \quad (24)$$

which will become useful in Section 3.2 when we analyze the long time properties of the fluctuations around the mean-field limit.

Remark 2.8. Assumptions 2.5 and 2.6 impose conditions on the initial measure μ_0 [12, 47, 51]. While the convergence of gradient flows in finite-dimensional Euclidean space to local minimizers is guaranteed under mild assumptions [39, 62], its infinite-dimensional counterpart, Assumption 2.6, may require further technical assumptions, left for future study. Also, while Assumption 2.5 implies that μ_∞ is a stationary point of (12), Assumption 2.6 does not imply that μ_∞ minimizes $\mathcal{L}$.

3 Fluctuations from Mean-Field Gradient Flow

The main goal of this section is to characterize the deviations of finite-particle shallow networks from their mean-field evolution, by first deriving an estimate for $f_t^{(m)} - f_t$ for $t \geq 0$ (Section 3.1), and then analyzing its long-time properties (Section 3.2). In Section 3.3, we then motivate a choice of the regularization term in (4) that controls the bound on the long-time fluctuations derived in Section 3.2, and which is also connected to generalization via the variation norm [5] or Barron norm [44] of functions.

3.1 A Dynamical Central Limit Theorem

Let us start by defining

$$g_t^{(m)} := m^{1/2}(f_t^{(m)} - f_t) . \quad (25)$$

By the static Central Limit Theorem (CLT) we know that, if we draw the initial values of the parameters θ_i independently from μ_0 as specified in Assumption 2.3, $g_0^{(m)}$ has a limit as $m \rightarrow \infty$, leading to estimates similar to (5) with $\mu^{(m)}$ and μ replaced by the initial $\mu_0^{(m)}$ and μ_0 , respectively. For $t > 0$, however, this estimate is not preserved by the gradient flow: the static CLT no longer applies and needs to be replaced by a dynamical variant [10, 60, 65, 66]. Next, we derive this dynamical CLT in the context of neural network optimization.

To this end let us define the discrepancy measure $\omega_t^{(m)}$ such that

$$\int_D \chi(\theta) \omega_t^{(m)}(d\theta) := m^{1/2} \int_D \chi(\theta) (\mu_t^{(m)}(d\theta) - \mu_t(d\theta)) , \quad (26)$$

for any continuous test function $\chi : D \rightarrow \mathbb{R}$. We can then represent $g_t^{(m)}$ in terms of $\omega_t^{(m)}$ as

$$g_t^{(m)} = \int_D \varphi(\theta, \cdot) \omega_t^{(m)}(d\theta) . \quad (27)$$

Hence, we will first establish how the limit of $\omega_t^{(m)}$ as $m \rightarrow \infty$ evolves over time. This can be done by noting that the representation formula (13) implies that

$$\int_D \chi(\theta) \omega_t^{(m)}(d\theta) = m^{1/2} \int_D (\chi(\Theta_t^{(m)}(\theta)) \mu_0^{(m)}(d\theta) - \chi(\Theta_t(\theta)) \mu_0(d\theta)) , \quad (28)$$

where $\Theta_t^{(m)}$ solves (15) with μ_0 replaced by $\mu_0^{(m)}$. Defining

$$\mathbf{T}_t^{(m)}(\theta) = m^{1/2}(\Theta_t^{(m)}(\theta) - \Theta_t(\theta)) , \quad (29)$$

we can write (28) as

$$\begin{aligned} \int_D \chi(\theta) \omega_t^{(m)}(d\theta) &= \int_D \chi(\Theta_t(\theta)) \omega_0^{(m)}(d\theta) \\ &\quad + \int_0^1 \int_D \nabla \chi(\Theta_t(\theta) + m^{-1/2} \eta \mathbf{T}_t^{(m)}(\theta)) \cdot \mathbf{T}_t^{(m)}(\theta) \mu_0^{(m)}(d\theta) d\eta . \end{aligned} \quad (30)$$

As shown in Appendix C.1, we can take the limit $m \rightarrow \infty$ of this formula to obtain:

Proposition 3.1 (Dynamical CLT - I). *Under Assumptions 2.2 and 2.3, $\forall t \geq 0$, as $m \rightarrow \infty$ we have $\omega_t^{(m)} \rightarrow \omega_t$ weakly in law with respect to $\mathbb{P}_0$, where ω_t is such that given a test function $\chi : D \rightarrow \mathbb{R}$,*

$$\int_D \chi(\boldsymbol{\theta}) \omega_t(d\boldsymbol{\theta}) = \int_D \chi(\boldsymbol{\Theta}_t(\boldsymbol{\theta})) \omega_0(d\boldsymbol{\theta}) + \int_D \nabla \chi(\boldsymbol{\Theta}_t(\boldsymbol{\theta})) \cdot \mathbf{T}_t(\boldsymbol{\theta}) \mu_0(d\boldsymbol{\theta}). \quad (31)$$

Here ω_0 is the Gaussian measure with mean zero and covariance

$$\mathbb{E}_0 [\omega_0(d\boldsymbol{\theta}) \omega_0(d\boldsymbol{\theta}')] = \mu_0(d\boldsymbol{\theta}) \delta_{\boldsymbol{\theta}}(d\boldsymbol{\theta}') - \mu_0(d\boldsymbol{\theta}) \mu_0(d\boldsymbol{\theta}'), \quad (32)$$

where $\mathbb{E}_0$ denotes expectation over $\mathbb{P}_0$, and $\mathbf{T}_t = \lim_{n \rightarrow \infty} m^{1/2}(\boldsymbol{\Theta}_t^{(m)} - \boldsymbol{\Theta}_t)$ is the flow solution to

$$\begin{aligned} \dot{\mathbf{T}}_t(\boldsymbol{\theta}) = & -\nabla \nabla V(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mu_t) \mathbf{T}_t(\boldsymbol{\theta}) - \int_D \nabla \nabla' K(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \boldsymbol{\Theta}_t(\boldsymbol{\theta}')) \mathbf{T}_t(\boldsymbol{\theta}') \mu_0(d\boldsymbol{\theta}') \\ & - \int_D \nabla K(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \boldsymbol{\Theta}_t(\boldsymbol{\theta}')) \omega_0(d\boldsymbol{\theta}') \end{aligned} \quad (33)$$

with initial condition $\mathbf{T}_0 = 0$ and where $\boldsymbol{\Theta}_t$ solves (14) and $\nabla \nabla V(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mu_t)$ is a shorthand for

$$\nabla \nabla V(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mu_t) = -\nabla \nabla F(\boldsymbol{\Theta}_t(\boldsymbol{\theta})) + \int_D \nabla \nabla K(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \boldsymbol{\Theta}_t(\boldsymbol{\theta}')) \mu_0(d\boldsymbol{\theta}'). \quad (34)$$

A direct consequence of this proposition and formula (27) is:

Corollary 3.2. *Under Assumptions 2.2 and 2.3, $\forall t \geq 0$, as $m \rightarrow \infty$ we have $g_t^{(m)} \rightarrow g_t$ pointwise in law with respect to $\mathbb{P}_0$, where g_t is given in terms of the limiting measure ω_t or the flow $\mathbf{T}_t$ as*

$$g_t = \int_D \varphi(\boldsymbol{\theta}, \cdot) \omega_t(d\boldsymbol{\theta}) = \int_D \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \cdot) \omega_0(d\boldsymbol{\theta}) + \int_D \nabla \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \cdot) \cdot \mathbf{T}_t(\boldsymbol{\theta}) \mu_0(d\boldsymbol{\theta}). \quad (35)$$

It is interesting to comment on the origin of both terms at the right hand side of (31) and, consequently, (35). The first term captures the deviations induced by fluctuations of $\mu_0^{(m)}$ around μ_0 assuming that the flow $\boldsymbol{\Theta}_t^{(m)}$ is unaffected by these fluctuations, and remains equal to $\boldsymbol{\Theta}_t$. In particular, this term is the one we would obtain if we were to resample $\mu_t^{(m)}$ from μ_t at every $t \geq 0$, i.e. use $\bar{\mu}_t^{(m)} = m^{-1} \sum_{i=1}^m \delta_{\bar{\boldsymbol{\theta}}_t^i}$ with $\{\bar{\boldsymbol{\theta}}_t^i\}_{i=1}^m$ sampled i.i.d. from μ_t , so that $\boldsymbol{\Theta}_t^{(m)}$ is identical to $\boldsymbol{\Theta}_t$ in (28). In this case, the limiting discrepancy measure $\bar{\omega}_t$ would simply be given by

$$\int_D \chi(\boldsymbol{\theta}) \bar{\omega}_t(d\boldsymbol{\theta}) = \int_D \chi(\boldsymbol{\Theta}_t(\boldsymbol{\theta})) \omega_0(d\boldsymbol{\theta}), \quad (36)$$

while the associated deviation in the represented function would read

$$\bar{g}_t = \int_D \varphi(\boldsymbol{\theta}, \cdot) \bar{\omega}_t(d\boldsymbol{\theta}) = \int_D \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \cdot) \omega_0(d\boldsymbol{\theta}). \quad (37)$$

The second term at right hand side of (31) and (35) captures the deviations to the flow $\boldsymbol{\Theta}_t$ in (15) induced by the perturbation of μ_0 , i.e. how much $\boldsymbol{\Theta}_t^{(m)}$ differs from $\boldsymbol{\Theta}_t$ in (28). In the limit as $m \rightarrow \infty$, these deviations are captured by the solution $\mathbf{T}_t$ to (33), as is apparent from (30).

The difference between g_t and $\bar{g}_t$ can also be quantified via the following Volterra equation, which can be derived from Proposition 3.1 and relates the evolution of g_t to that of $\bar{g}_t$.

Corollary 3.3 (Dynamical CLT - II). *Under Assumptions 2.2 and 2.3, $\forall t \geq 0$, pointwise on Ω , we have $g_t^{(m)} \rightarrow g_t$ in law with respect to $\mathbb{P}_0$ as $m \rightarrow \infty$, where g_t solves the Volterra equation*

$$g_t(\mathbf{x}) + \int_0^t \int_{\Omega} \Gamma_{t,s}(\mathbf{x}, \mathbf{x}') g_s(\mathbf{x}') \hat{\nu}(d\mathbf{x}') ds = \bar{g}_t(\mathbf{x}) . \quad (38)$$

Here $\bar{g}_t$ is given in (37) and we defined

$$\Gamma_{t,s}(\mathbf{x}, \mathbf{x}') = \int_D \langle \nabla_{\theta} \varphi(\Theta_t(\theta)), J_{t,s}(\theta) \nabla_{\theta} \varphi(\Theta_s(\theta)) \rangle \mu_0(d\theta) , \quad (39)$$

where $J_{t,s}$ is the solution to

$$\frac{d}{dt} J_{t,s}(\theta) = -\nabla \nabla V(\Theta_t(\theta), \mu_t) J_{t,s}(\theta), \quad J_{s,s}(\theta) = Id . \quad (40)$$

This corollary is proven in Appendix C.2. In a nutshell, (38) can be established using Duhamel's principle on (33) by considering all terms at the right hand side except the first as the source term (hence the role of $J_{t,s}$) and inserting the result in (35).

3.2 Long-Time Behavior of the Fluctuations

Next, we study the long-time behavior of g_t and, in particular, evaluate

$$\lim_{t \rightarrow \infty} \mathbb{E}_0 \|g_t\|_{\hat{\nu}}^2 = \lim_{t \rightarrow \infty} \lim_{m \rightarrow \infty} m \mathbb{E}_0 \|f_t^{(m)} - f_t\|_{\hat{\nu}}^2 . \quad (41)$$

This limit quantifies the asymptotic approximation error of $f_t^{(m)}$ around its mean field limit f_t after gradient flow, i.e. if we take $m \rightarrow \infty$ first, then $t \rightarrow \infty$ – taking these limits in the opposite order is of interest too but is beyond the scope of the present paper. Our main result is to show that, under certain assumptions to be specified below, the limit in (41) is not only finite but necessarily upper-bounded by $\lim_{t \rightarrow \infty} \mathbb{E}_0 \|\bar{g}_t\|_{\hat{\nu}}^2$ with $\bar{g}_t$ given in (37). That is, the approximation error at the end of training is always no higher than that obtained by resampling the mean-field measure μ_{∞} defined in Proposition 2.7.

It is useful to start by considering an idealized case, namely when the initial conditions are sampled as in Assumption 2.3 with $\mu_0 = \mu_{\infty}$. In that case, there is no evolution at mean field level, i.e. $\Theta_t(\theta) = \Theta_{\infty}(\theta) = \theta$, $\mu_t = \mu_{\infty}$, and $f_t = f_{\infty} = \int_D \varphi_{\infty}(\theta, \cdot) \mu_{\infty}(d\theta)$, but the CLT fluctuations still evolve. In particular, it is easy to see that the Volterra equation in (38) for g_t becomes

$$g_t(\mathbf{x}) + \int_0^t \int_{\Omega} \Gamma_{t-s}^{\infty}(\mathbf{x}, \mathbf{x}') g_s(\mathbf{x}') \hat{\nu}(d\mathbf{x}') ds = \bar{g}_{\infty}(\mathbf{x}) . \quad (42)$$

Here $\Gamma_{t-s}^{\infty}(\mathbf{x}, \mathbf{x}')$ is the Volterra kernel obtained by solving (40) with $\nabla \nabla V(\Theta_t(\theta), \mu_t)$ replaced by $\nabla \nabla V(\theta, \mu_{\infty})$ and inserting the result in (39) with $\Theta_t(\theta) = \theta$ and $\mu_0 = \mu_{\infty}$,

$$\Gamma_{t-s}^{\infty}(\mathbf{x}, \mathbf{x}') = \int_D \langle \nabla_{\theta} \varphi(\theta, \mathbf{x}), e^{-(t-s)\nabla \nabla V(\theta, \mu_{\infty})} \nabla_{\theta} \varphi(\theta, \mathbf{x}') \rangle \mu_{\infty}(d\theta) , \quad (43)$$

and $\bar{g}_{\infty}$ is the Gaussian field with variance

$$\mathbb{E}_0 \|\bar{g}_{\infty}\|_{\hat{\nu}}^2 = \int_D \|\varphi(\theta, \cdot)\|_{\hat{\nu}}^2 \mu_{\infty}(d\theta) - \|f_{\infty}\|_{\hat{\nu}}^2 . \quad (44)$$

From (23) in Proposition 2.7 we know that $\nabla \nabla V(\boldsymbol{\theta}, \mu_\infty)$ is positive semidefinite for μ_∞ -almost all $\boldsymbol{\theta}$. As a result, we prove in D.1 that the Volterra kernel (43) viewed as an operator on functions defined on $\Omega \times [0, T]$ is positive semidefinite. Therefore, we have

$$\begin{aligned} \int_0^T \|g_t\|_{\hat{\nu}}^2 dt &\leq \int_0^T \|g_t\|_{\hat{\nu}}^2 dt + \int_0^T \int_0^t \int_{\Omega \times \Omega} g_t(\mathbf{x}) \Gamma_{t-s}^\infty(\mathbf{x}, \mathbf{x}') g_s(\mathbf{x}') \hat{\nu}(d\mathbf{x}) \hat{\nu}(d\mathbf{x}') ds dt \\ &= \int_0^T \mathbb{E}_{\hat{\nu}}(g_t \bar{g}_\infty) dt \leq T^{1/2} \|\bar{g}_\infty\|_{\hat{\nu}} \left(\int_0^T \|g_t\|_{\hat{\nu}}^2 dt \right)^{1/2}. \end{aligned} \quad (45)$$

Together with (44), this implies that

Theorem 3.4. *Under Assumptions 2.2, 2.3, 2.5 and 2.6, with $\mu_0 = \mu_\infty$ and μ_∞ as specified in Proposition 2.7, we have*

$$\lim_{T \rightarrow \infty} \frac{1}{T} \int_0^T \mathbb{E}_0 \|g_t\|_{\hat{\nu}}^2 dt \leq \int_D \|\varphi(\boldsymbol{\theta}, \cdot)\|_{\hat{\nu}}^2 \mu_\infty(d\boldsymbol{\theta}) - \|f_\infty\|_{\hat{\nu}}^2. \quad (46)$$

This theorem indicates that, if we knew μ_∞ and could sample initial conditions for the parameters from it, it would still be favorable to train these parameters as this would reduce the approximation error. Of course, in practice we have no *a priori* access to μ_∞ , and so the relevant question is whether (46) also holds if we sample initial conditions from any μ_0 such that Proposition 2.7 holds.

In light of (35), one way to address this question is to study the long-time behavior of T_t . In the setup without regularization ($\lambda = 0$), we can do so by leveraging existing results that, under certain assumptions, the mean-field gradient flow converges to a global minimizer which interpolates the training data points exactly [12, 47, 52, 60]. In this case, the following theorem shows that we can actually obtain stronger controls on the fluctuations than (46), which we prove in Appendix D.2.

Theorem 3.5 (Long-time fluctuations in the unregularized case). *Consider the ERM setting with $\lambda = 0$ and under Assumptions 2.2, 2.3 and 2.5. Suppose that as $t \rightarrow \infty$, μ_t converges to a global minimizer μ_∞ that interpolates the data, i.e. the function $f_\infty = \int_D \varphi(\boldsymbol{\theta}, \cdot) \mu_\infty(d\boldsymbol{\theta})$ satisfies*

$$\forall \mathbf{x} \in \text{supp } \hat{\nu} : f_\infty(\mathbf{x}) = f_*(\mathbf{x}), \quad (47)$$

and, furthermore, the convergence satisfies

$$\int_0^\infty |\mathcal{L}(\mu_t)|^{1/2} dt < \infty \quad (48)$$

Then (46) holds. Additionally,

1. if Assumption 2.1 also holds, i.e., in the shallow neural network setting, we further have

$$\lim_{T \rightarrow \infty} \frac{1}{T} \int_0^T \mathbb{E}_0 \|g_t\|_{\hat{\nu}}^2 dt = 0; \quad (49)$$

2. if $\mu_0 = \mu_\infty$, then $\|g_t\|_{\hat{\nu}}$ decreases monotonically in t .

Hence, in the shallow neural networks setting and under these assumptions, the fluctuations will eventually vanish in the $O(m^{-1/2})$ scale of CLT. Note that for (48) to hold, it is sufficient that $\mathcal{L}(\mu_t)$ decays at an asymptotic rate of $O(t^{-\alpha})$ with $\alpha > 2$. We leave the search for weaker sufficient conditions for future work.

When the limiting measure μ_∞ does not necessarily interpolate the training data, such as when regularization is added, we can proceed with the analysis of the long-time behavior of T_t under the following assumption on the long-time behavior of the curvature:

Theorem 3.6 (Long-time fluctuations under assumptions on the curvature). *Let $\Lambda_t(\boldsymbol{\theta})$ denote the smallest eigenvalue of the tensor $\nabla\nabla V(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mu_t)$ defined in (34) and assume that for a constant C (to be specified in Appendix D.3) such that*

$$-\int_D \min\{\Lambda_t(\boldsymbol{\theta}), 0\} \mu_0(d\boldsymbol{\theta}) = O(e^{-Ct}) \quad \text{as } t \rightarrow \infty. \quad (50)$$

Then (46) holds.

This theorem is proven in Appendix D.3. To intuitively understand (50), note that we know from (23) in Proposition 2.7 that $\Lambda_t(\boldsymbol{\theta}) \rightarrow 0$ μ_0 -almost surely as $t \rightarrow \infty$. Condition (50) can therefore be satisfied by having $\Lambda_t(\boldsymbol{\theta})$ converge to zero sufficiently fast in the regions of D where it is negative, or having the measure of these regions with respect to μ_0 converge to zero sufficiently fast, or both.

Alternatively, in the regularized ($\lambda > 0$) ERM setting, we can obtain the following result when the support of μ_∞ is atomic, as expected on general grounds [5, 9, 18, 26, 71]:

Theorem 3.7 (Long-time fluctuations in the regularized case). *Consider the ERM setting under Assumptions 2.2, 2.3 and 2.5. Suppose further that as $t \rightarrow \infty$, μ_t converges to μ_∞ satisfying*

$$\exists \sigma > 0 \text{ s.t. } \forall \boldsymbol{\theta} \in \text{supp } \mu_\infty : \nabla\nabla V(\boldsymbol{\theta}, \mu_\infty) \succ \sigma \text{Id}, \text{ and} \quad (51)$$

$$\boldsymbol{\Theta}_t \text{ admits an asymptotic uniform convergence rate of } O(t^{-\alpha}) \text{ with } \alpha > 3/2. \quad (52)$$

Then (46) holds with the “lim” replaced by “lim sup” on its LHS.

Theorem 3.7 is proven in Appendix D.4 by analyzing directly the Volterra equation (38) and establishing that its solution coincides with that of (42) in the limit as $t \rightarrow \infty$, a property that we also expect to hold more generally than under the assumptions of Theorem 3.7. In fact, we prove in Appendix D.4 that (52) can be replaced by a weaker condition, (237). We also discuss the relation between Theorem 3.7 and the work of [11] in Appendix D.4.3.

3.3 The Monte-Carlo Bound and Regularization

The bound (46) on the long-time fluctuations motivates us to control the term $\int_D \|\varphi(\boldsymbol{\theta}, \cdot)\|_{\hat{\nu}}^2 \mu_\infty(d\boldsymbol{\theta})$ using a suitable choice of regularization in (4). In the following, we restrict our attention to the shallow neural networks setting, and further assume that

Assumption 3.8. $\hat{D}$ is compact.

Under this assumption, there is

$$\int_D \|\varphi(\boldsymbol{\theta}, \cdot)\|_{\hat{\nu}}^2 \mu(d\boldsymbol{\theta}) = \int_D \int_\Omega |\varphi(\boldsymbol{\theta}, \mathbf{x})|^2 \hat{\nu}(d\mathbf{x}) \mu(d\boldsymbol{\theta}) \leq \hat{K}_M \int_D c^2 \mu(d\boldsymbol{\theta}), \quad (53)$$

where $\hat{K}_M = \max_{\mathbf{z} \in \hat{D}} \|\hat{\varphi}(\mathbf{z}, \cdot)\|_{\hat{\nu}}^2$. Thus, we consider regularization with $r(\boldsymbol{\theta}) = \frac{1}{2}c^2$, in which case (4) becomes

$$\min_{\mu \in \mathcal{P}(D)} \mathcal{L}(\mu) \quad \text{with} \quad \mathcal{L}(\mu) := \frac{1}{2} \|f[\mu] - f_*\|_{\hat{\nu}}^2 + \frac{1}{2}\lambda \int_D c^2 \mu(d\boldsymbol{\theta}). \quad (54)$$

Interestingly, this choice of regularization leads to learning in the function space $\mathcal{F}_1$ [5] (or alternatively, the Barron space [44]) associated with $\hat{\varphi}$, which is equipped with the *variation norm* (or the *Barron norm*) defined as

$$|\gamma_q(f)| := \inf_{\mu \in \mathcal{P}(D)} \left\{ \int_D |c|^q \mu(d\boldsymbol{\theta}); f(\mathbf{x}) = \int_D c \hat{\varphi}(\mathbf{z}, \mathbf{x}) \mu(d\boldsymbol{\theta}) \right\} = |\gamma_1(f)|^q, \quad q \geq 1. \quad (55)$$

We call $\int_D |c|^q \mu(d\theta)$ the q -norm of μ . One can verify [44, Proposition 1] that indeed, using any $q \geq 1$ above yields the same norm because μ , the object defining the integral representation (2), is in fact a *lifted* version of a more ‘fundamental’ object $\gamma = \int_{\mathbb{R}} c\mu(dc, \cdot) \in \mathcal{M}(\hat{D})$, the space of signed Radon measures over $\hat{D}$. They are related via the projection

$$\int_{\hat{D}} \chi(z) \gamma(dz) = \int_D c\chi(z) \mu(d\theta) \quad (56)$$

for all continuous test functions $\chi : \hat{D} \rightarrow \mathbb{R}$. One can also verify [11] that $\gamma_1(f) = \inf\{\|\gamma\|_{\text{TV}}; f(x) = \int_{\hat{D}} \hat{\varphi}(z, x) \gamma(dz)\}$, where $\|\gamma\|_{\text{TV}}$ is the *total variation* of γ [5].

The space $\mathcal{F}_1$ contains any RKHS whose kernel is generated as an expectation over features $k(x, x') = \int_{\hat{D}} \hat{\varphi}(z, x) \hat{\varphi}(z, x') \hat{\mu}_0(dz)$ with a base measure $\hat{\mu}_0 \in \mathcal{P}(\hat{D})$, but it provides crucial approximation advantages over such RKHS at approximating certain non-smooth, high-dimensional functions with hidden low-dimensional structure, giving rise to powerful generalization guarantees [5]. This also motivates the study of overparametrized shallow networks with the scaling as in (1), as opposed to the NTK scaling of $m^{-1/2}$ [36].

To learn in $\mathcal{F}_1$, a canonical approach is to consider the ERM problem

$$\min_{f \in \mathcal{F}_1} \frac{1}{2} \|f - f_*\|_{\hat{\nu}}^2 + \frac{1}{2} \lambda \gamma_1(f), \quad (57)$$

By (55), this is indeed equivalent to (54). In Appendix E, we prove the following proposition, which shows that the measure obtained from (54) indeed has its 2-norm controlled:

Proposition 3.9. *Under Assumptions 2.1, 2.2, and 3.8, $\mathcal{L}$ has no local minima and its global minimum value can only be attained at measures $\mu_\lambda \in \mathcal{P}(D)$ such that both $f_\lambda = \int_D \varphi(\theta, \cdot) \mu_\lambda(d\theta)$ and $c_\lambda = \int_D |c| \mu_\lambda(d\theta) = (\int_D |c|^2 \mu_\lambda(d\theta))^{1/2} \leq \gamma_1(f_*)$ are unique, and*

$$\lambda^2 |c_\lambda|^2 \hat{K}_M^{-1} \leq \|f_\lambda - f_*\|_{\hat{\nu}}^2, \quad \|f_\lambda - f_*\|_{\hat{\nu}}^2 + \lambda |c_\lambda|^2 \leq \lambda |\gamma_1(f_*)|^2. \quad (58)$$

where $\hat{K}_M = \max_{z \in \hat{D}} \|\hat{\varphi}(z, \cdot)\|_{\hat{\nu}}^2$.

4 Numerical Experiments

4.1 Student-Teacher Setting

We first perform numerical experiments in a student-teacher setting, using a shallow teacher network as the target function to be learned by shallow student networks with different widths m of the hidden layer. Both $\hat{D}$ and Ω are taken to be the unit sphere of $d = 16$ dimensions, and we take $\hat{\varphi}(z, x) = \max(0, \langle z, x \rangle)$. The teacher network has two neurons, (c_1, z_1) and (c_2, z_2) , in the hidden layer, with $c_1 = c_2 = 1$ and z_1 and z_2 sampled i.i.d. from the uniform distribution on $\hat{D}$ and then fixed across the experiments. We vary the width of the student network in the range of $m = 128, 256, 512, 1024$ and 2048 , with their initial z_i ’s sampled i.i.d. from the uniform distribution on $\hat{D}$. We consider two ways for initializing the c_i ’s of the student networks: 1) *Gaussian-initialization*, where the c_i ’s are sampled i.i.d. from $\mathcal{N}(0, 1)$; and 2) *zero-initialization*, where each c_i is set to be 0.

We train the student networks in two ways: using the *population* loss or the *empirical* loss. For the former scenario, the data distribution ν is chosen to be uniform on Ω , which allows an analytical formula for the loss as well as its gradient. The student networks are trained by gradient descent under L_2 loss. Moreover, we rescale both the squared loss and the gradient by d in order to adjust to the $\frac{1}{d}$ factor resulting from spherical integrals, and set the learning rate (which is the step size for

discretizing (9)) to be 1. The models are trained for 20000 epochs. For each choice of m , we run the experiment $\kappa = 20$ times with different random initializations of the student network. The average fluctuation of the population loss is defined as $\frac{1}{\kappa} \sum_{k=1}^{\kappa} \|f_k^{(m)} - \bar{f}^{(m)}\|_{\nu}^2$ for the population loss, with $\bar{f}^{(m)} = \frac{1}{\kappa} \sum_{k=1}^{\kappa} f_k^{(m)}$ being the averaged model, similar to the approach in [28]. The other plotted quantities – loss, TV-norm and 2-norm – are averaged across the κ number of runs. The TV-norm (i.e., 1-norm) and 2-norm are defined as in Appendix 3.3.

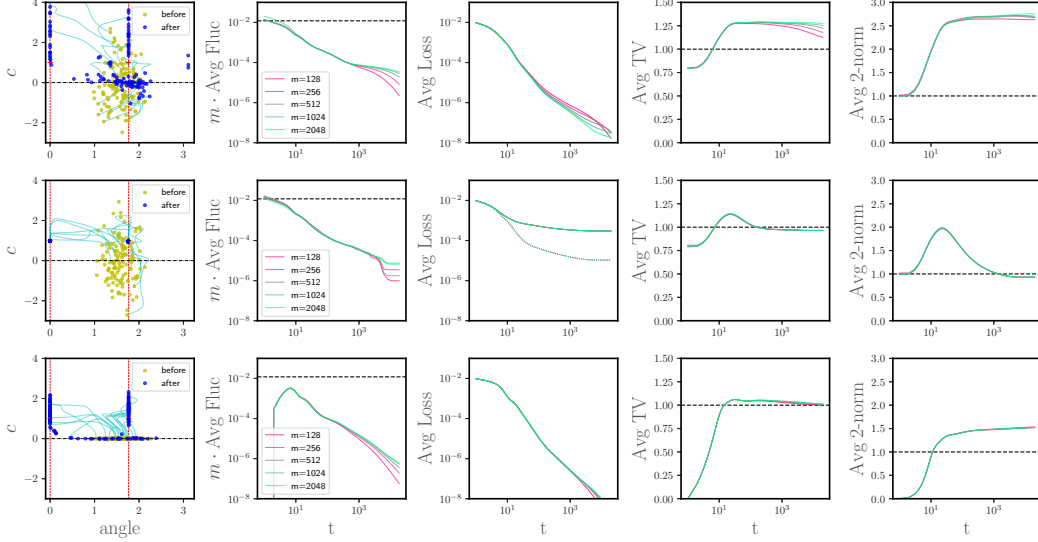


Figure 1: Results of the experiments in the student-teacher setting and where the student networks are trained by gradient descent on the *population* loss. Each row corresponds to one setup. *Row 1*: Using unregularized loss and non-zero-initialization; *Row 2*: Using regularized loss with $\lambda = 0.01$ and non-zero-initialization; *Row 3*: Using unregularized loss and zero-initialization. In each row, *Column 1* plots the trajectory of the neurons, $\theta_i = (c_i, z_i)$, of a student network of width 128 during its training, with x -coordinate being the angle between z_i and that of a chosen teacher’s neuron and y -coordinate being c_i . The yellow dots, blue dots and cyan curves mark their initial values, terminal values, and trajectory during training. *Columns 2-5* plot the average fluctuations (scaled by m), average loss, average TV norm, and average 2-norm during training, respectively, computed across $\kappa = 20$ runs with different random initializations of the student network for each choice of m . In *Column 2*, the *solid* curves give the average fluctuation of the population loss and the black horizontal *dashed* line gives an approximate value of the asymptotic Monte-Carlo bound in (46) for this setting computed in Appendix F. In *Column 3*, the *solid* curves indicate the total population loss, and the *dotted* curves indicate the unregularized population loss (for the regularized case only). In *Columns 4 and 5*, the horizontal *dashed* line gives the relevant norm of the teacher network.

The results for the scenario of training under the *population* loss are presented in Figures 1. As seen from *Column 3* the average loss values remain similar over time for different choices of m , justifying the approximation by a mean-field dynamics. In the unregularized case with non-zero initialization, the fluctuation of the population loss (shown in *Column 2*) remains close to a $1/m$ scaling in roughly the first 10^3 epochs, after which it decays faster for smaller m . Interestingly, this coincides with the tendency for the student neurons with z not aligned with the teacher neurons to slowly have their $|c|$ decrease to zero due to a finite- m effect, which is also reflected in the decrease in TV-norm. Aside from this phenomenon, the fluctuations decay at similar rates for different choices of m , which is consistent with our theory, since their dynamics are governed by the same dynamical CLT. Also, when regularization is added, each student neuron becomes aligned with one of the teacher neurons in both z and c after training; without regularization but using zero-initialization,

after training, each student neuron either becomes aligned with one of the teacher neurons in z or has c close to zero. Both of these choices result in lower TV-norms and 2-norms compared to using non-zero initialization and without regularization.

Next, we consider the *empirical* loss scenario (ERM setting), using $n = 32$ random vectors sampled i.i.d. from the uniform distribution ν on Ω as the training dataset, which then define the empirical data measure $\hat{\nu}(dx) = \frac{1}{n} \sum_{l=1}^n \delta_{x_l}(dx)$. We use the full training dataset for computing the gradient at every iteration. The other training setups are the same as when the population loss is used. We additionally plot the average fluctuation of the training loss, defined as $\frac{1}{\kappa} \sum_{k=1}^{\kappa} \|f_k^{(m)} - \bar{f}^{(m)}\|_{\hat{\nu}}^2$.

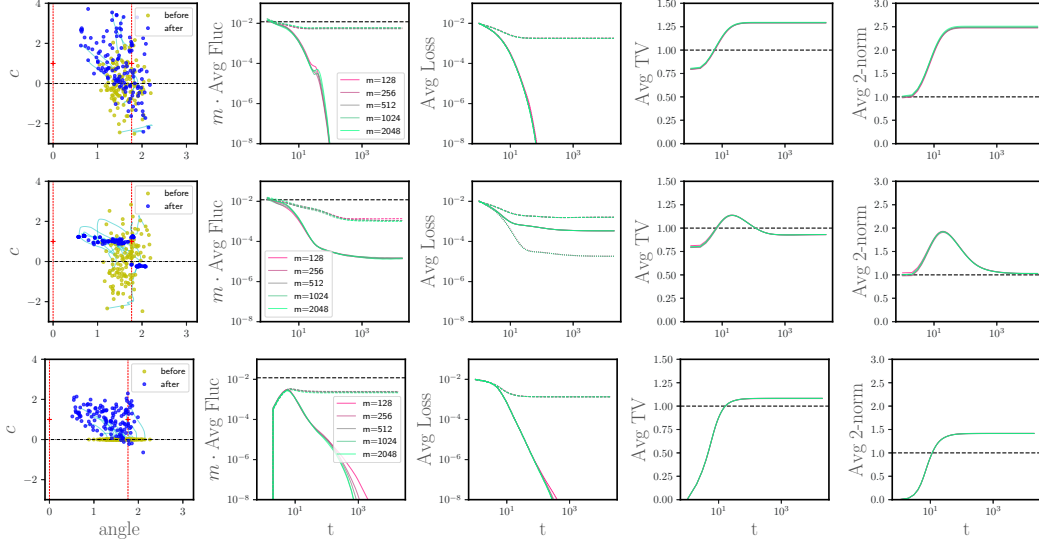


Figure 2: Results of the experiments in the student-teacher setting and where the student networks are trained by gradient descent on the *empirical* loss. In Column 2, the *solid* curves indicate the average fluctuation in the *training* loss, the *dashed* curves indicate the average fluctuation in the *population* loss computed analytically via spherical integrals, and the black horizontal *dashed* line indicates an approximate value of the asymptotic Monte-Carlo bound in (46) for this setting computed in Appendix F. In Column 3, the *solid* curves indicate the total *training* loss, the *dotted* curves the unregularized *training* loss (for the regularized case only), and the *dashed* curves the unregularized *population* loss. All the other plot settings are identical to Figure 1.

The results for the scenario of training under the *empirical* loss is presented in Figures 2. Compared to the scenario of training under the *population* loss, we see that in the unregularized cases, both the average training loss and the average fluctuation of the training loss decay to below 10^{-8} within 10^3 iterations, and the latter observation is consistent with (49). In the regularized case, neither of them vanishes, but the average fluctuation of the training loss indeed remains below the asymptotic Monte-Carlo bound given in (46), whose analytical expression and numerical value in this setup (under the approximation of replacing μ_∞ , f_∞ and $\hat{\nu}$ by the target measure, the target function and ν , respectively) are given in Appendix F. Regularization and zero-initialization have a weaker effect in aligning the student neurons with the teacher neurons after training compared to the scenario of training under the *population* loss, but they still result in lower TV-norm and 2-norm, and moreover, lower average fluctuation and (slightly) lower average value of the *population* loss. This demonstrates their positive effects on both approximation and generalization.

4.2 Non-planted Case

We also conducted experiments in which the target function is not given by a teacher network but rather by $f_*(\mathbf{x}) = \int_{\hat{D}} \hat{\varphi}(z, \mathbf{x}) \hat{\mu}_*(dz)$, where $\hat{\mu}_*$ is the uniform measure on the 1-dimensional great circle in the first 2 dimensions, i.e., $\{(\cos \theta, \sin \theta), 0, \dots, 0 : \theta \in [0, 2\pi)\} \subseteq \mathbb{S}^d$, and where $\hat{D}, \Omega, \hat{\varphi}$ as well as the widths of the student networks remain the same as in the previous experiments. The student networks are trained using gradient descent under the population loss where the data distribution ν is uniform on Ω , which allows an analytical formula for the gradient using spherical integrals.

The results are shown in Figure 3. We observe that the behaviors of the fluctuation are qualitatively similar to those found in Figure 1.

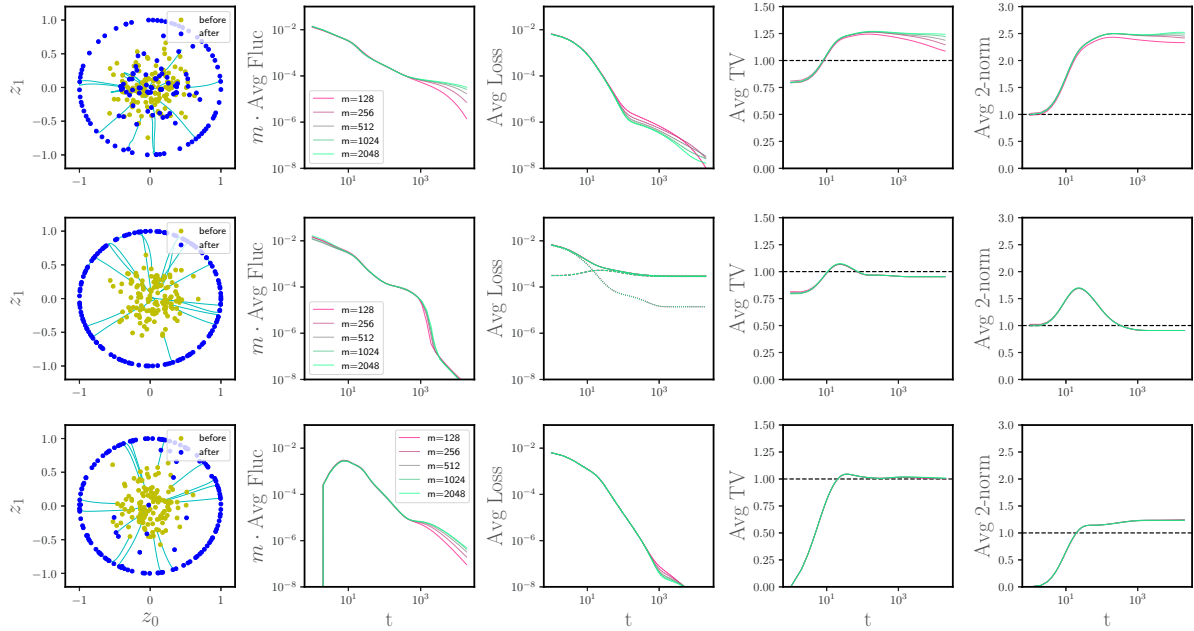


Figure 3: Results of the experiments with a non-planted target using the exact population loss, as described in Section 4.2. *Row 1:* Using unregularized loss and non-zero-initialization; *Row 2:* Using regularized loss with $\lambda = 0.01$ and non-zero-initialization; *Row 3:* Using unregularized loss and zero-initialization. In each row, *Column 1* plots the projection in the first two dimensions of the z_i 's in the student network. The other columns show the same quantities as in Figure 1.

5 Conclusions

Here we studied the deviations of shallow neural networks from their infinite-width limit, and how these deviations evolve during training by gradient flow. In the ERM setting, we established that under different sets of conditions, the long-term deviation under the Central Limit Theorem (CLT) scaling is controlled by a Monte Carlo (MC) resampling error, giving width-asymptotic guarantees that do not depend on the data dimension explicitly. The MC resampling bound motivates a choice of regularization that is also connected to generalization via the variation-norm function spaces.

Our results thus seem to paint a favorable picture for high-dimensional learning, in which the optimization and generalization guarantees for the idealized mean-field limit could be transferred to their finite-width counterparts. However, we stress that these results are asymptotic, in that we take limits both in the width and time. In the face of negative results for the computational

efficiency of training shallow networks [20, 31, 41, 45, 54], an important challenge is to leverage additional structure in the problem (such as the empirical data distribution [32], or the structure of the minimizers [18]) to provide nonasymptotic versions of our results, along the lines of [11] or [40]. Finally, another clear direction for future research is to extend our techniques to deep neural architectures, in light of recent works that consider deep or residual models [3, 25, 42, 49, 59, 68].

Acknowledgements

This work benefited from discussions with Lénia Chizat and Carles Domingo-Enrich. Z.C. acknowledges support from the Henry MacCracken Fellowship. G.M.R. acknowledges support from the James S. McDonnell Foundation. J.B. acknowledges support from the Alfred P. Sloan Foundation, NSF RI-1816753, NSF CAREER CIF 1845360, and the Institute for Advanced Study. E. V.-E. acknowledges support from the National Science Foundation (NSF) Materials Research Science and Engineering Center Program Grant No. DMR-1420073, and from NSF Grant No. DMS-1522767.

References

- [1] Ben Adlam and Jeffrey Pennington. The neural tangent kernel in high dimensions: Triple descent and a multi-scale theory of generalization.
- [2] Zeyuan Allen-Zhu, Yuanzhi Li, and Zhao Song. A convergence theory for deep learning via over-parameterization. In *International Conference on Machine Learning*, pages 242–252, 2019.
- [3] Dyego Araújo, Roberto I Oliveira, and Daniel Yukimura. A mean-field limit for certain deep neural networks. *arXiv preprint arXiv:1906.00193*, 2019.
- [4] Sanjeev Arora, Simon S Du, Wei Hu, Zhiyuan Li, and Ruosong Wang. Fine-grained analysis of optimization and generalization for overparameterized two-layer neural networks. *arXiv preprint arXiv:1901.08584*, 2019.
- [5] Francis Bach. Breaking the curse of dimensionality with convex neural networks. *The Journal of Machine Learning Research*, 18(1):629–681, 2017.
- [6] Yu Bai and Jason D Lee. Beyond linearization: On quadratic and higher-order approximation of wide neural networks. *arXiv preprint arXiv:1910.01619*, 2019.
- [7] Javier Baladron, Diego Fasoli, Olivier Faugeras, and Jonathan Touboul. Mean field description of and propagation of chaos in recurrent multipopulation networks of hodgkin-huxley and fitzhugh-nagumo neurons. *arXiv preprint arXiv:1110.4294*, 2011.
- [8] Mikhail Belkin, Daniel Hsu, Siyuan Ma, and Soumik Mandal. Reconciling modern machine learning practice and the bias-variance trade-off. *arXiv:1812.11118 [cs, stat]*, September 2019. arXiv: 1812.11118.
- [9] Claire Boyer, Antonin Chambolle, Yohann De Castro, Vincent Duval, Frédéric De Gournay, and Pierre Weiss. On representer theorems and convex regularization. *SIAM Journal on Optimization*, 29(2):1260–1281, 2019.
- [10] Werner Braun and K Hepp. The vlasov dynamics and its fluctuations in the $1/n$ limit of interacting classical particles. *Communications in mathematical physics*, 56(2):101–113, 1977.

- [11] Lenaïc Chizat. Sparse optimization on measures with over-parameterized gradient descent. *arXiv preprint arXiv:1907.10300*, 2019.
- [12] Lenaïc Chizat and Francis Bach. On the global convergence of gradient descent for over-parameterized models using optimal transport. In *Advances in Neural Information Processing Systems*, pages 3036–3046, 2018.
- [13] Lénaïc Chizat and Francis Bach. Implicit bias of gradient descent for wide two-layer neural networks trained with the logistic loss. *arXiv preprint arXiv:2002.04486*, 2020.
- [14] Lenaïc Chizat, Edouard Oyallon, and Francis Bach. On lazy training in differentiable programming. In *Advances in Neural Information Processing Systems*, pages 2937–2947, 2019.
- [15] Youngmin Cho and Lawrence K. Saul. Kernel methods for deep learning. In Y. Bengio, D. Schuurmans, J. D. Lafferty, C. K. I. Williams, and A. Culotta, editors, *Advances in Neural Information Processing Systems 22*, pages 342–350. Curran Associates, Inc., 2009.
- [16] Roberto Cortez. Uniform propagation of chaos for kac’s 1d particle system. *Journal of Statistical Physics*, 165(6):1102–1113, 2016.
- [17] Valentin De Bortoli, Alain Durmus, Xavier Fontaine, and Umut Simsekli. Quantitative propagation of chaos for sgd in wide neural networks. *arXiv preprint arXiv:2007.06352*, 2020.
- [18] Jaume de Dios and Joan Bruna. On sparsity in overparametrised shallow relu networks. *arXiv preprint arXiv:2006.10225*, 2020.
- [19] Pierre Del Moral and Laurent Miclo. Branching and interacting particle systems approximations of feynman-kac formulae with applications to non-linear filtering. In *Seminaire de probabilites XXXIV*, pages 1–145. Springer, 2000.
- [20] Ilias Diakonikolas, Daniel M Kane, Vasilis Kontonis, and Nikos Zarifis. Algorithms and sq lower bounds for pac learning one-hidden-layer relu networks. In *Conference on Learning Theory*, pages 1514–1539, 2020.
- [21] Simon S Du, Jason D Lee, Haochuan Li, Liwei Wang, and Xiyu Zhai. Gradient descent finds global minima of deep neural networks. *arXiv preprint arXiv:1811.03804*, 2018.
- [22] Alain Durmus, Andreas Eberle, Arnaud Guillin, and Raphael Zimmer. An elementary approach to uniform in time propagation of chaos. *arXiv preprint arXiv:1805.11387*, 2018.
- [23] Ethan Dyer and Guy Gur-Ari. Asymptotics of wide networks from feynman diagrams. *arXiv preprint arXiv:1909.11304*, 2019.
- [24] Weinan E, Chao Ma, and Lei Wu. Machine learning from a continuous viewpoint, 2019.
- [25] Cong Fang, Jason D Lee, Pengkun Yang, and Tong Zhang. Modeling from features: a mean-field framework for over-parameterized deep neural networks. *arXiv preprint arXiv:2007.01452*, 2020.
- [26] SD Fisher and Joseph W Jerome. Spline solutions to l1 extremal problems in one and several variables. *Journal of Approximation Theory*, 13(1):73–83, 1975.

- [27] Mario Geiger, Arthur Jacot, Stefano Spigler, Franck Gabriel, Levent Sagun, Stéphane d’Ascoli, Giulio Biroli, Clément Hongler, and Matthieu Wyart. Scaling description of generalization with number of parameters in deep learning. *Journal of Statistical Mechanics: Theory and Experiment*, 2020(2):023401, 2020.
- [28] Mario Geiger, Stefano Spigler, Arthur Jacot, and Matthieu Wyart. Disentangling feature and lazy learning in deep neural networks: an empirical study. *arXiv preprint arXiv:1906.08034*, 2019.
- [29] Behrooz Ghorbani, Song Mei, Theodor Misiakiewicz, and Andrea Montanari. Limitations of lazy training of two-layers neural network. In *Advances in Neural Information Processing Systems*, pages 9111–9121, 2019.
- [30] Behrooz Ghorbani, Song Mei, Theodor Misiakiewicz, and Andrea Montanari. When do neural networks outperform kernel methods? *arXiv preprint arXiv:2006.13409*, 2020.
- [31] Surbhi Goel, Aravind Gollakota, Zhihan Jin, Sushrut Karmalkar, and Adam Klivans. Super-polynomial lower bounds for learning one-layer neural networks using gradient descent. *arXiv preprint arXiv:2006.12011*, 2020.
- [32] Sebastian Goldt, Galen Reeves, Marc Mézard, Florent Krzakala, and Lenka Zdeborová. The gaussian equivalence of generative models for learning with two-layer neural networks. *arXiv preprint arXiv:2006.14709*, 2020.
- [33] G. Gripenberg, S. O. Londen, and O. Staffans. *Volterra Integral and Functional Equations*. Encyclopedia of Mathematics and its Applications. Cambridge University Press, 1990.
- [34] Boris Hanin and Mihai Nica. Finite depth and width corrections to the neural tangent kernel. *arXiv preprint arXiv:1909.05989*, 2019.
- [35] Jiaoyang Huang and Horng-Tzer Yau. Dynamics of deep neural networks and neural tangent hierarchy. *arXiv preprint arXiv:1909.08156*, 2019.
- [36] Arthur Jacot, Franck Gabriel, and Clément Hongler. Neural tangent kernel: Convergence and generalization in neural networks. In *Advances in neural information processing systems*, pages 8571–8580, 2018.
- [37] Carlo Lancellotti. On the fluctuations about the vlasov limit for n-particle systems with mean-field interactions. *Journal of Statistical Physics*, 136(4):643–665, 2009.
- [38] Jaehoon Lee, Yasaman Bahri, Roman Novak, Samuel S Schoenholz, Jeffrey Pennington, and Jascha Sohl-Dickstein. Deep neural networks as gaussian processes. *arXiv preprint arXiv:1711.00165*, 2017.
- [39] Jason D Lee, Ioannis Panageas, Georgios Piliouras, Max Simchowitz, Michael I Jordan, and Benjamin Recht. First-order methods almost always avoid saddle points. *arXiv preprint arXiv:1710.07406*, 2017.
- [40] Yuanzhi Li, Tengyu Ma, and Hongyang R. Zhang. Learning over-parametrized two-layer neural networks beyond ntk. volume 125 of *Proceedings of Machine Learning Research*, pages 2613–2682. PMLR, 09–12 Jul 2020.

- [41] Roi Livni, Shai Shalev-Shwartz, and Ohad Shamir. On the computational efficiency of training neural networks. In *Advances in Neural Information Processing Systems*, pages 855–863, 2014.
- [42] Yiping Lu, Chao Ma, Yulong Lu, Jianfeng Lu, and Lexing Ying. A mean-field analysis of deep resnet and beyond: Towards provable optimization via overparameterization from depth. *arXiv preprint arXiv:2003.05508*, 2020.
- [43] Tao Luo, Zhi-Qin John Xu, Zheng Ma, and Yaoyu Zhang. Phase diagram for two-layer relu neural networks at infinite-width limit. *arXiv preprint arXiv:2007.07497*, 2020.
- [44] Chao Ma, Lei Wu, and Weinan E. Barron spaces and the compositional function spaces for neural network models. *arXiv preprint arXiv:1906.08039*, 2019.
- [45] Pasin Manurangsi and Daniel Reichman. The computational complexity of training relu (s). *arXiv preprint arXiv:1810.04207*, 2018.
- [46] Song Mei, Theodor Misiakiewicz, and Andrea Montanari. Mean-field theory of two-layers neural networks: dimension-free bounds and kernel limit. *arXiv preprint arXiv:1902.06015*, 2019.
- [47] 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*, 115(33):E7665–E7671, 2018.
- [48] Brady Neal, Sarthak Mittal, Aristide Baratin, Vinayak Tantia, Matthew Scicluna, Simon Lacoste-Julien, and Ioannis Mitliagkas. A modern take on the bias-variance tradeoff in neural networks. *arXiv preprint arXiv:1810.08591*, 2018.
- [49] Phan-Minh Nguyen and Huy Tuan Pham. A rigorous framework for the mean field limit of multilayer neural networks. *arXiv preprint arXiv:2001.11443*, 2020.
- [50] Grant Rotskoff, Samy Jelassi, Joan Bruna, and Eric Vanden-Eijnden. Global convergence of neuron birth-death dynamics. *arXiv preprint arXiv:1902.01843*, 2019.
- [51] Grant Rotskoff and Eric Vanden-Eijnden. Parameters as interacting particles: long time convergence and asymptotic error scaling of neural networks. In *Advances in Neural Information Processing Systems*, pages 7146–7155, 2018.
- [52] 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. *arXiv preprint arXiv:1805.00915*, 2018.
- [53] Nicolas Le Roux and Yoshua Bengio. Continuous neural networks. volume 2 of *Proceedings of Machine Learning Research*, pages 404–411, San Juan, Puerto Rico, 21–24 Mar 2007. PMLR.
- [54] Itay Safran and Ohad Shamir. Spurious local minima are common in two-layer ReLU neural networks. In *International Conference on Machine Learning*, pages 4433–4441, 2018.
- [55] Samir Salem. A gradient flow approach of uniform in time propagation of chaos for particles in double a well confinement. *arXiv preprint arXiv:1810.08946*, 2018.
- [56] Jamil Salhi, James MacLaurin, and Salwa Toumi. On uniform propagation of chaos. *Stochastics*, 90(1):49–60, 2018.

- [57] Sylvia Serfaty. Coulomb gases and ginzburg-landau vortices. *arXiv preprint arXiv:1403.6860*, 2014.
- [58] Justin Sirignano and Konstantinos Spiliopoulos. Dgm: A deep learning algorithm for solving partial differential equations. *Journal of Computational Physics*, 375:1339–1364, 2018.
- [59] Justin Sirignano and Konstantinos Spiliopoulos. Mean field analysis of deep neural networks. *arXiv preprint arXiv:1903.04440*, 2019.
- [60] Justin Sirignano and Konstantinos Spiliopoulos. Mean field analysis of neural networks: A central limit theorem. *Stochastic Processes and their Applications*, 130(3):1820–1852, 2020.
- [61] Justin Sirignano and Konstantinos Spiliopoulos. Mean field analysis of neural networks: A law of large numbers. *SIAM Journal on Applied Mathematics*, 80(2):725–752, 2020.
- [62] Stephen Smale. Stable manifolds for differential equations and diffeomorphisms. *Annali della Scuola Normale Superiore di Pisa-Classe di Scienze*, 17(1-2):97–116, 1963.
- [63] 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*, 65(2):742–769, 2018.
- [64] Stefano Spigler, Mario Geiger, Stéphane d’Ascoli, Levent Sagun, Giulio Biroli, and Matthieu Wyart. A jamming transition from under-to over-parametrization affects loss landscape and generalization. *arXiv preprint arXiv:1810.09665*, 2018.
- [65] Herbert Spohn. *Large scale dynamics of interacting particles*. Springer Science & Business Media, 2012.
- [66] Alain-Sol Sznitman. Topics in propagation of chaos. In *Ecole d’été de probabilités de Saint-Flour XIX—1989*, pages 165–251. Springer, 1991.
- [67] Luca Venturi, Afonso S. Bandeira, and Joan Bruna. Spurious valleys in one-hidden-layer neural network optimization landscapes. *Journal of Machine Learning Research*, 20(133):1–34, 2019.
- [68] Stephan Wojtowytsch et al. On the banach spaces associated with multi-layer relu networks: Function representation, approximation theory and gradient descent dynamics. *arXiv preprint arXiv:2007.15623*, 2020.
- [69] Blake Woodworth, Suriya Gunasekar, Jason D Lee, Edward Moroshko, Pedro Savarese, Itay Golan, Daniel Soudry, and Nathan Srebro. Kernel and rich regimes in overparametrized models. *arXiv preprint arXiv:2002.09277*, 2020.
- [70] Chiyuan Zhang, Samy Bengio, Moritz Hardt, Benjamin Recht, and Oriol Vinyals. Understanding deep learning requires rethinking generalization. *arXiv preprint arXiv:1611.03530*, 2016.
- [71] S. Zuhovickii. Remarks on problems in approximation theory. *Mat. Zbirnik KDU*, 1948.

Appendix

Table of Contents

A	Notations	21
B	Long-Time Properties of the Mean-Field Gradient Flow	22
C	Derivations of the Dynamical Central Limit Theorem	24
C.1	Proof of Proposition 3.1 (Dynamical CLT - I)	24
C.2	Proof of Proposition 3.3 (Dynamical CLT - II)	26
D	Long-Time Behavior of the Fluctuations	26
D.1	Proof of Theorem 3.4 ($\mu_0 = \mu_\infty$ case)	26
D.2	Proof of Theorem 3.5 (Unregularized case)	27
D.3	Proof of Theorem 3.6 (Under assumptions on the curvature in the long-time)	41
D.4	Proof of Theorem 3.7 (Regularized case)	44
E	Properties of the Minimizers of the Regularized Loss	55
F	Analytical Calculations of the Resampling Error	58

A Notations

We will use $\nabla\varphi(\theta, x)$ and $\nabla\nabla\varphi(\theta, x)$ to denote $\nabla_\theta\varphi(\theta, x)$ and $\nabla_\theta\nabla_\theta\varphi(\theta, x)$, respectively. We will use $\nabla K(\theta, \theta')$ to denote $\nabla_\theta K(\theta, \theta')$, $\nabla\nabla K(\theta, \theta')$ to denote $\nabla_\theta\nabla_\theta K(\theta, \theta')$, $\nabla'\nabla K(\theta, \theta')$ to denote $\nabla_{\theta'}\nabla_\theta K(\theta, \theta')$, and $\nabla'\nabla'K(\theta, \theta')$ to denote $\nabla_{\theta'}\nabla_{\theta'}K(\theta, \theta')$. We will write $V_t(\cdot)$ for $V(\cdot, \mu_t)$ and $V_\infty(\cdot)$ for $V(\cdot, \mu_\infty)$.

Let $D' = \cup_{t>0} \text{supp } \mu_t$. Under Assumption 2.5 and Proposition 2.7, D' is bounded, and we denote its diameter by $|D'|$. We will use C_φ , $C_{\nabla\varphi}$ and $C_{\nabla\nabla\varphi}$ to denote the supremum of $|\varphi(\theta, x)|$, $|\nabla\varphi(\theta, x)|$ and $|\nabla\nabla\varphi(\theta, x)|$ over $\theta \in D'$ and $x \in \text{supp } \hat{\nu}$, which are all finite under Assumptions 2.2 and the boundedness of D' . We will use $L_{\nabla\nabla\varphi}$ to denote the (uniform-in- x) Lipschitz constant of $\nabla\nabla\varphi(\theta, x)$ in θ , which is also finite under Assumption 2.2.

The following notations will be used in Appendix D.2: Assuming that D is Euclidean (under Assumption 2.2), let $\mathcal{V}(D)$ denote the space of random vector fields on D . It becomes a Hilbert space once equipped with the inner product

$$\langle \xi_1, \xi_2 \rangle_0 := \mathbb{E}_0 \int_D \xi_1(\theta) \cdot \xi_2(\theta) \mu_0(d\theta), \quad (59)$$

where ξ_1, ξ_2 denotes two random vector fields in $\mathcal{V}(D)$. This inner product gives rise to the norm

$$\|\xi\|_0^2 := \mathbb{E}_0 \int_D |\xi(\theta)|^2 \mu_0(d\theta). \quad (60)$$

For each t , we define $\mathbf{b}_t \in \mathcal{V}(D)$ as

$$\mathbf{b}_t(\boldsymbol{\theta}) = \int_D \nabla K(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \boldsymbol{\Theta}_t(\boldsymbol{\theta}')) \omega_0(d\boldsymbol{\theta}') \quad (61)$$

which depends on the random measure ω_0 . We define two linear operators, $\mathcal{A}_t^{(K)}$ and $\mathcal{A}_t^{(V)}$ on $\mathcal{V}(D)$, as

$$(\mathcal{A}_t^{(K)} \boldsymbol{\xi})(\boldsymbol{\theta}) = \int_D \nabla' \nabla K(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \boldsymbol{\Theta}_t(\boldsymbol{\theta}')) \boldsymbol{\xi}(\boldsymbol{\theta}') \mu_0(d\boldsymbol{\theta}') \quad (62)$$

$$= \int_{\Omega} \nabla \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mathbf{x}) \left(\int_D \nabla \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}'), \mathbf{x})^\top \boldsymbol{\xi}(\boldsymbol{\theta}') \mu_0(d\boldsymbol{\theta}') \right) \hat{\nu}(d\mathbf{x}), \quad (63)$$

$$(\mathcal{A}_t^{(V)} \boldsymbol{\xi})(\boldsymbol{\theta}) = \nabla \nabla V(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mu_t) \boldsymbol{\xi}(\boldsymbol{\theta}), \quad (64)$$

for $\boldsymbol{\xi} \in \mathcal{V}(D)$. Under Assumption 2.5, we also define $\mathbf{b}_\infty$, $\mathcal{A}_\infty^{(K)}$, and $\mathcal{A}_\infty^{(V)}$ similarly by replacing $\boldsymbol{\Theta}_t(\cdot)$ with $\boldsymbol{\Theta}_\infty(\cdot)$.

Let $\mathcal{W}_n(\Omega)$ denote the space of random functions on Ω . It becomes a Hilbert space once equipped with the inner product

$$\langle \eta_1, \eta_2 \rangle_{\hat{\nu}, 0} := \mathbb{E}_0 \int_{\Omega} \eta_1(\mathbf{x}) \eta_2(\mathbf{x}) \hat{\nu}(d\mathbf{x}) = \frac{1}{n} \mathbb{E}_0 \sum_{l=1}^n \eta_1(\mathbf{x}_l) \eta_2(\mathbf{x}_l), \quad (65)$$

which gives rise to the norm

$$\|\eta\|_{\hat{\nu}, 0}^2 := \langle \eta, \eta \rangle_{\hat{\nu}, 0} = \mathbb{E}_0 \|\eta\|_{\hat{\nu}}^2. \quad (66)$$

With an abuse of notation, we will consider elements in $\mathcal{W}_n(\Omega)$ equivalently as random vectors on $\mathbb{R}^L$. Next, we can define $\mathcal{B}_t$ to be the operator that maps $\eta \in \mathcal{W}_n(\Omega)$ into the vector field

$$(\mathcal{B}_t \eta)(\boldsymbol{\theta}) = \int_{\Omega} \nabla \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mathbf{x}) \eta(\mathbf{x}) \hat{\nu}(d\mathbf{x}) \quad (67)$$

in $\mathcal{V}(D)$. Its transpose is

$$(\mathcal{B}_t^\top \boldsymbol{\xi})(\mathbf{x}) = \int_D \nabla \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mathbf{x}) \boldsymbol{\xi}(\boldsymbol{\theta}) \mu_0(d\boldsymbol{\theta}), \quad (68)$$

which maps a vector field $\boldsymbol{\xi} \in \mathcal{V}(D)$ back into $\mathcal{W}_n(\Omega)$.

B Long-Time Properties of the Mean-Field Gradient Flow

Proof of Proposition 2.7: The compactness of $\cup_{t \geq 0} \text{supp } \mu_t$ follows from (20) and the compactness of $\text{supp } \mu_0$ assumed in Assumption 2.3. $\mu_t \rightharpoonup \mu_\infty$ follows from (13) and (20).

Under Assumption 2.5, $\boldsymbol{\Theta}_\infty$ is a local minimizer of the energy $\mathcal{E}$ defined in (18). Consider a local perturbation $\epsilon \boldsymbol{\Theta}_\Delta$ to $\boldsymbol{\Theta}$. The energy value after the perturbation is

$$\begin{aligned} \mathcal{E}(\boldsymbol{\Theta}_\infty + \epsilon \boldsymbol{\Theta}_\Delta) &= - \int_D F(\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}) + \epsilon \boldsymbol{\Theta}_\Delta(\boldsymbol{\theta})) \mu_0(d\boldsymbol{\theta}) \\ &\quad + \frac{1}{2} \int_D \int_D K(\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}) + \epsilon \boldsymbol{\Theta}_\Delta(\boldsymbol{\theta}), \boldsymbol{\Theta}_\infty(\boldsymbol{\theta}') + \epsilon \boldsymbol{\Theta}_\Delta(\boldsymbol{\theta}')) \mu_0(d\boldsymbol{\theta}') \mu_0(d\boldsymbol{\theta}). \end{aligned} \quad (69)$$

Under Assumptions 2.2, using Taylor expansion, we have

$$\begin{aligned} F(\Theta_\infty(\theta) + \epsilon \Theta_\Delta(\theta)) &= F(\Theta_\infty(\theta)) + \epsilon \nabla F(\Theta_\infty(\theta)) \cdot \Theta_\Delta(\theta) \\ &\quad + \frac{1}{2} \epsilon^2 \langle \Theta_\Delta(\theta), \nabla \nabla F(\Theta_\infty(\theta)) \Theta_\Delta(\theta) \rangle + O(\epsilon^3) \end{aligned} \quad (70)$$

$$\begin{aligned} &K(\Theta_\infty(\theta) + \epsilon \Theta_\Delta(\theta), \Theta_\infty(\theta') + \epsilon \Theta_\Delta(\theta')) \\ &= K(\Theta_\infty(\theta), \Theta_\infty(\theta')) + \epsilon \nabla K(\Theta_\infty(\theta), \Theta_\infty(\theta')) \Theta_\Delta(\theta) \\ &\quad + \epsilon \nabla' K(\Theta_\infty(\theta), \Theta_\infty(\theta')) \Theta_\Delta(\theta') + \frac{1}{2} \epsilon^2 \langle \Theta_\Delta(\theta), \nabla \nabla K(\Theta_\infty(\theta), \Theta_\infty(\theta')) \Theta_\Delta(\theta) \rangle \\ &\quad + \frac{1}{2} \epsilon^2 \langle \Theta_\Delta(\theta'), \nabla' \nabla' K(\Theta_\infty(\theta), \Theta_\infty(\theta')) \Theta_\Delta(\theta') \rangle \\ &\quad + \epsilon^2 \langle \Theta_\Delta(\theta), \nabla' \nabla K(\Theta_\infty(\theta), \Theta_\infty(\theta')) \Theta_\Delta(\theta') \rangle + O(\epsilon^3). \end{aligned} \quad (71)$$

Hence, there is

$$\begin{aligned} &\mathcal{E}(\Theta_\infty + \epsilon \Theta_\Delta) - \mathcal{E}(\Theta_\infty) \\ &= \epsilon \int_D \left(-\nabla F(\Theta_\infty(\theta)) + \int_D \nabla K(\Theta_\infty(\theta), \Theta_\infty(\theta')) \mu_0(d\theta') \right) \Theta_\Delta(\theta) \mu_0(d\theta) \\ &\quad + \frac{1}{2} \epsilon^2 \left(\int_D \langle \Theta_\Delta(\theta), \left(\nabla \nabla F(\Theta_\infty(\theta)) + \int_D \nabla \nabla K(\Theta_\infty(\theta), \Theta_\infty(\theta')) \mu_0(d\theta') \right) \Theta_\Delta(\theta) \rangle \mu_0(d\theta) \right. \\ &\quad \left. + \int_D \int_D \langle \Theta_\Delta(\theta), \nabla' \nabla K(\Theta_\infty(\theta), \Theta_\infty(\theta')) \Theta_\Delta(\theta') \rangle \mu_0(d\theta) \mu_0(d\theta') \right) + O(\epsilon^3). \end{aligned} \quad (72)$$

Since Θ_Δ is arbitrary can ϵ can be taken arbitrarily small, we see that for Θ_∞ to be a local minimizer, the first-order condition is, $\forall \theta \in \text{supp } \mu_0$,

$$-\nabla F(\Theta_\infty(\theta)) + \int_D \nabla K(\Theta_\infty(\theta), \Theta_\infty(\theta')) \mu_0(d\theta') = 0, \quad (73)$$

or

$$\nabla V(\Theta_\infty(\theta), \mu_\infty) = 0, \quad (74)$$

and the second-order condition is, $\forall \Theta_\Delta$,

$$\begin{aligned} &\int_D \langle \Theta_\Delta(\theta), \left(\nabla \nabla F(\Theta_\infty(\theta)) + \int_D \nabla \nabla K(\Theta_\infty(\theta), \Theta_\infty(\theta')) \mu_0(d\theta') \right) \Theta_\Delta(\theta) \rangle \mu_0(d\theta) \\ &\quad + \int_D \int_D \langle \Theta_\Delta(\theta), \nabla' \nabla K(\Theta_\infty(\theta), \Theta_\infty(\theta')) \Theta_\Delta(\theta') \rangle \mu_0(d\theta) \mu_0(d\theta') \geq 0, \end{aligned} \quad (75)$$

or

$$\begin{aligned} &\int_D \langle \Theta_\Delta(\theta), \nabla \nabla V(\Theta_\infty(\theta), \mu_\infty) \Theta_\Delta(\theta) \rangle \mu_0(d\theta) \\ &\quad + \int_D \int_D \langle \Theta_\Delta(\theta), \nabla' \nabla K(\Theta_\infty(\theta), \Theta_\infty(\theta')) \Theta_\Delta(\theta') \rangle \mu_0(d\theta) \mu_0(d\theta') \geq 0. \end{aligned} \quad (76)$$

Suppose for contradiction that $\exists D^- \subseteq D$ with $\mu_0(D^-) > 0$ such that $\nabla \nabla V(\Theta_\infty(\theta), \mu_\infty)$ is not positive semidefinite. Define $\Lambda_\infty(\theta)$ to be the least eigenvalue of $\nabla \nabla V(\Theta_\infty(\theta), \mu_\infty)$. Then there is $\Lambda_\infty(\theta) < 0$ on D^- . In addition, $\exists \zeta > 0$, $\exists D_0^- \subseteq D^-$ with $\mu_0(D_0^-) > 0$ such that $\Lambda_\infty(\theta) < -\zeta$. For $\theta \in D_0^-$, let $\Theta_{\Delta,0}(\theta)$ be a normalized eigenvector to $\nabla \nabla V(\Theta_\infty(\theta), \mu_\infty)$ associated with its least

eigenvalue. Moreover, for $J \in \mathbb{N}^*$ that is large enough, we can select any subset $D_J^- \subset D_0^-$ such that $\mu_0(D_J^-) = \frac{1}{J} < \mu_0(D_0^-)$. Then, define

$$\Theta_{\Delta,J}(\theta) = J^{1/2} \mathbb{1}_{\theta \in D_J^-} \Theta_{\Delta,0}(\theta), \quad (77)$$

Then, there is

$$\begin{aligned} & \int_D \int_D \langle \Theta_{\Delta}(\theta), \nabla' \nabla K(\Theta_{\infty}(\theta), \Theta_{\infty}(\theta')) \Theta_{\Delta}(\theta') \rangle \mu_0(d\theta) \mu_0(d\theta') \\ &= \int_{\Omega} \left| \int_D \nabla \varphi(\Theta_{\infty}(\theta), x) \Theta_{\Delta,n} \mu_0(d\theta) \right|^2 \hat{\nu}(dx) \\ &= \int_{\Omega} \left| J^{1/2} \int_{D_J^-} \nabla \varphi(\Theta_{\infty}(\theta), x) \Theta_{\Delta,0} \mu_0(d\theta) \right|^2 \hat{\nu}(dx) \\ &\leq C_{\nabla \varphi}^2 J^{-1}. \end{aligned} \quad (78)$$

On the other hand

$$\begin{aligned} & \int_D \langle \Theta_{\Delta,J}(\theta), \nabla \nabla V(\Theta_{\infty}(\theta), \mu_{\infty}) \Theta_{\Delta,J}(\theta) \rangle \mu_0(d\theta) \\ &= \int_{D_J^-} J^{-1} \langle \Theta_{\Delta,0}(\theta), \nabla \nabla V(\Theta_{\infty}(\theta), \mu_{\infty}) \Theta_{\Delta,0}(\theta) \rangle \mu_0(d\theta) \\ &\leq -\zeta. \end{aligned} \quad (79)$$

Therefore, for J large enough, we will have

$$\begin{aligned} & \int_D \int_D \langle \Theta_{\Delta}(\theta), \nabla' \nabla K(\Theta_{\infty}(\theta), \Theta_{\infty}(\theta')) \Theta_{\Delta}(\theta') \rangle \mu_0(d\theta) \mu_0(d\theta') \\ &+ \int_D \langle \Theta_{\Delta,n}(\theta), \nabla \nabla V(\Theta_{\infty}(\theta), \mu_{\infty}) \Theta_{\Delta,J}(\theta) \rangle \mu_0(d\theta) < 0, \end{aligned} \quad (80)$$

which contradicts (76). Hence, we can conclude that μ_0 -almost surely, $\nabla \nabla V(\Theta_{\infty}(\theta), \mu_{\infty})$ is positive semidefinite.

C Derivations of the Dynamical Central Limit Theorem

C.1 Proof of Proposition 3.1 (Dynamical CLT - I)

The following derivation is an adaptation of the approach in [10] for Vlasov interacting particle systems to our scenario. To start, Θ_t and $\Theta_t^{(m)}$ are governed by the following equations, respectively:

$$\begin{aligned} \dot{\Theta}_t(\theta) &= -\nabla V(\Theta_t(\theta), \mu_t), & \Theta_0(\theta) &= \theta \\ \dot{\Theta}_t^{(m)}(\theta) &= -\nabla V(\Theta_t^{(m)}(\theta), \mu_t^{(m)}), & \Theta_0^{(m)}(\theta) &= \theta \end{aligned} \quad (81)$$

Taking the difference between the two equations in (81) and using the mean value theorem, we get

$$\begin{aligned}
& \dot{\mathbf{T}}_t^{(m)}(\boldsymbol{\theta}) \\
&= m^{1/2} \left(\dot{\boldsymbol{\Theta}}_t^{(m)}(\boldsymbol{\theta}) - \dot{\boldsymbol{\Theta}}_t(\boldsymbol{\theta}) \right) \\
&= -m^{1/2} \left(\nabla V(\boldsymbol{\Theta}_t^{(m)}(\boldsymbol{\theta}), \mu_t^{(m)}) - \nabla V(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mu_t) \right) \\
&= -m^{1/2} \left(\nabla V(\boldsymbol{\Theta}_t^{(m)}(\boldsymbol{\theta}), \mu_t) - \nabla V(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mu_t) \right) - m^{1/2} \left(\nabla V(\boldsymbol{\Theta}_t, \mu_t^{(m)}) - \nabla V(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mu_t) \right) \\
&\quad - m^{1/2} \left[\left(\nabla V(\boldsymbol{\Theta}_t^{(m)}(\boldsymbol{\theta}), \mu_t^{(m)}) - \nabla V(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mu_t^{(m)}) \right) - \left(\nabla V(\boldsymbol{\Theta}_t^{(m)}, \mu_t) - \nabla V(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mu_t) \right) \right] \\
&= -\nabla \nabla V(\tilde{\boldsymbol{\Theta}}_{t,1}^{(m)}(\boldsymbol{\theta}), \mu_t) \mathbf{T}_t^{(m)}(\boldsymbol{\theta}) - \int_D \nabla K(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \boldsymbol{\theta}') \omega_t^{(m)}(d\boldsymbol{\theta}') \\
&\quad - m^{-1/2} \left(\int_D \nabla \nabla K(\tilde{\boldsymbol{\Theta}}_{t,2}^{(m)}(\boldsymbol{\theta}), \boldsymbol{\theta}') \omega_t^{(m)}(d\boldsymbol{\theta}') \right) \mathbf{T}_t^{(m)}(\boldsymbol{\theta}),
\end{aligned} \tag{82}$$

where $\tilde{\boldsymbol{\Theta}}_{t,1}^{(m)}(\boldsymbol{\theta})$ and $\tilde{\boldsymbol{\Theta}}_{t,2}^{(m)}(\boldsymbol{\theta})$ denote points that lie on the line segment between $\boldsymbol{\Theta}_t(\boldsymbol{\theta})$ and $\boldsymbol{\Theta}_t^{(m)}(\boldsymbol{\theta})$. Using (30), we can substitute $\omega_t^{(m)}$ in the second term at the right hand side, for which we get

$$\begin{aligned}
\int_D \nabla K(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \boldsymbol{\theta}') \omega_t^{(m)}(d\boldsymbol{\theta}') &= \int_D \nabla K(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \boldsymbol{\Theta}_t(\boldsymbol{\theta}')) \omega_0^{(m)}(d\boldsymbol{\theta}') \\
&\quad + \int_D \nabla' \nabla K(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \tilde{\boldsymbol{\Theta}}_{t,3}^{(m)}(\boldsymbol{\theta}')) \mathbf{T}_t^{(m)}(\boldsymbol{\theta}') \mu_0(d\boldsymbol{\theta}') \\
&\quad + m^{-1/2} \int_D \nabla' \nabla K(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \tilde{\boldsymbol{\Theta}}_{t,3}^{(m)}(\boldsymbol{\theta}')) \mathbf{T}_t^{(m)}(\boldsymbol{\theta}') \omega_0^{(m)}(d\boldsymbol{\theta}').
\end{aligned} \tag{83}$$

Therefore, under Assumption 2.2, we have

$$\begin{aligned}
\dot{\mathbf{T}}_t^{(m)}(\boldsymbol{\theta}) &= -\nabla \nabla V(\tilde{\boldsymbol{\Theta}}_{t,1}^{(m)}, \mu_t) \mathbf{T}_t^{(m)}(\boldsymbol{\theta}) \\
&\quad - \int_D \nabla' \nabla K(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \tilde{\boldsymbol{\Theta}}_{t,3}^{(m)}(\boldsymbol{\theta}')) \mathbf{T}_t^{(m)}(\boldsymbol{\theta}') \mu_0(d\boldsymbol{\theta}') \\
&\quad - \int_D \nabla K(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \boldsymbol{\Theta}_t(\boldsymbol{\theta}')) \omega_0^{(m)}(d\boldsymbol{\theta}') + O(m^{-1/2}).
\end{aligned} \tag{84}$$

Now, we consider the limit as $m \rightarrow \infty$. By the standard CLT, we have that $\omega_0^{(m)}(d\boldsymbol{\theta}) \rightarrow \omega_0(d\boldsymbol{\theta})$ weakly with respect to $\mathbb{P}_0$, where $\omega_0(d\boldsymbol{\theta})$ is the Gaussian measure with mean zero and covariance defined in (32). On the other hand, by finite-time LLN, we have $\boldsymbol{\Theta}_t^{(m)}(\boldsymbol{\theta}) \rightarrow \boldsymbol{\Theta}_t(\boldsymbol{\theta})$ pointwise, $\mathbb{P}_0$ -almost surely, and as a consequence $\tilde{\boldsymbol{\Theta}}_{t,1}^{(m)}(\boldsymbol{\theta}), \tilde{\boldsymbol{\Theta}}_{t,3}^{(m)}(\boldsymbol{\theta}) \rightarrow \boldsymbol{\Theta}_t(\boldsymbol{\theta})$ as well. Therefore, $\mathbf{T}_t^{(m)}(\boldsymbol{\theta}) \rightarrow \mathbf{T}_t(\boldsymbol{\theta})$ pointwise, $\mathbb{P}_0$ -almost surely, where the limiting $\mathbf{T}_t(\boldsymbol{\theta})$ solves the equation obtained by taking the limit $m \rightarrow \infty$ on both sides of (84), which becomes (33). (33) should be solved with initial condition $\mathbf{T}_0(\boldsymbol{\theta}) = 0$ since $\mathbf{T}_0^{(m)}(\boldsymbol{\theta}) = m^{1/2}(\boldsymbol{\Theta}_0^{(m)}(\boldsymbol{\theta}) - \boldsymbol{\Theta}_0(\boldsymbol{\theta})) = 0$.

Finally, taking the limit $m \rightarrow \infty$ on both sides of the equation (30), we deduce that $\omega_t^{(m)}(d\boldsymbol{\theta}) \rightarrow \omega_t(d\boldsymbol{\theta})$ weakly, in law with respect to $\mathbb{P}_0$, where the limiting $\omega_t(d\boldsymbol{\theta})$ satisfies

$$\int_D \chi(\boldsymbol{\theta}) \omega_t(d\boldsymbol{\theta}) = \int_D \chi(\boldsymbol{\Theta}_t(\boldsymbol{\theta})) \omega_0(d\boldsymbol{\theta}) + \int_D \nabla \chi(\boldsymbol{\Theta}_t(\boldsymbol{\theta})) \cdot \mathbf{T}_t(\boldsymbol{\theta}) \mu_0(d\boldsymbol{\theta}). \tag{85}$$

This ends the proof of Proposition 3.1. $\square$

C.2 Proof of Proposition 3.3 (Dynamical CLT - II)

Recall from (33) that

$$\begin{aligned}\dot{\mathbf{T}}_t(\boldsymbol{\theta}) &= -\nabla \nabla V(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mu_t) \mathbf{T}_t(\boldsymbol{\theta}) - \int_D \nabla' \nabla K(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \boldsymbol{\Theta}_t(\boldsymbol{\theta}')) \mathbf{T}_t(\boldsymbol{\theta}') \mu_0(d\boldsymbol{\theta}') \\ &\quad - \int_D \nabla K(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \boldsymbol{\Theta}_t(\boldsymbol{\theta}')) \omega_0(d\boldsymbol{\theta}') \\ &= -\nabla \nabla V(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mu_t) \mathbf{T}_t(\boldsymbol{\theta}) - \int_D \nabla K(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \boldsymbol{\theta}') \omega_t(d\boldsymbol{\theta}').\end{aligned}\tag{86}$$

Since $\mathbf{T}_0(\boldsymbol{\theta}) = 0$, we can use Duhamel's principle to deduce that

$$\begin{aligned}\mathbf{T}_t(\boldsymbol{\theta}) &= - \int_0^t J_{t,s}(\boldsymbol{\theta}) \int_D \nabla K(\boldsymbol{\Theta}_s(\boldsymbol{\theta}), \boldsymbol{\theta}') \omega_s(d\boldsymbol{\theta}') ds \\ &= - \int_0^t \int_\Omega J_{t,s}(\boldsymbol{\theta}) \nabla \varphi(\boldsymbol{\Theta}_s(\boldsymbol{\theta}), \mathbf{x}) \int_D \varphi(\boldsymbol{\theta}', \mathbf{x}) \omega_s(d\boldsymbol{\theta}') \hat{\nu}(d\mathbf{x}) ds \\ &= - \int_0^t \int_\Omega J_{t,s}(\boldsymbol{\theta}) \nabla \varphi(\boldsymbol{\Theta}_s(\boldsymbol{\theta}), \mathbf{x}) g_s(\mathbf{x}) \hat{\nu}(d\mathbf{x}) ds,\end{aligned}\tag{87}$$

where the tensor $J_{t,s}(\boldsymbol{\theta})$ is the Jacobian defined in Proposition 3.3. As a result

$$\begin{aligned}g_t(\mathbf{x}) &= \int_D \varphi(\boldsymbol{\theta}, \mathbf{x}) \omega_t(d\boldsymbol{\theta}) \\ &= \int_D \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mathbf{x}) \omega_0(d\boldsymbol{\theta}) + \int_D \nabla \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mathbf{x}) \cdot \mathbf{T}_t(\boldsymbol{\theta}) \mu_0(d\boldsymbol{\theta}) \\ &= \int_D \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mathbf{x}) \omega_0(d\boldsymbol{\theta}) \\ &\quad - \int_D \int_0^t \int_\Omega \langle \nabla \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mathbf{x}), J_{t,s}(\boldsymbol{\theta}) \nabla \varphi(\boldsymbol{\Theta}_s(\boldsymbol{\theta}), \mathbf{x}') \rangle g_s(\mathbf{x}') \hat{\nu}(d\mathbf{x}') ds \mu_0(d\boldsymbol{\theta}) \\ &= \bar{g}_t(\mathbf{x}) - \int_0^t \int_\Omega \int_D \langle \nabla \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mathbf{x}), J_{t,s}(\boldsymbol{\theta}) \nabla \varphi(\boldsymbol{\Theta}_s(\boldsymbol{\theta}), \mathbf{x}') \rangle \mu_0(d\boldsymbol{\theta}) g_s(\mathbf{x}') \hat{\nu}(d\mathbf{x}') ds \\ &= \bar{g}_t(\mathbf{x}) - \int_0^t \int_\Omega \Gamma_{t,s}(\mathbf{x}, \mathbf{x}') g_s(\mathbf{x}') \hat{\nu}(d\mathbf{x}') ds,\end{aligned}\tag{88}$$

with $\bar{g}_t(\mathbf{x})$ and $\Gamma_{t,s}(\mathbf{x}, \mathbf{x}')$ defined in (37) and (39), respectively. This is (38). $\square$

D Long-Time Behavior of the Fluctuations

D.1 Proof of Theorem 3.4 ($\mu_0 = \mu_\infty$ case)

With the argument outlined in Section 3.2, what remains to be shown is that Γ_{t-s}^∞ is positive-semidefinite as a Volterra kernel, according to the definition in [33]. We will utilize the following known result:

Proposition D.1 (Gripenberg et al. [33]). *Let $k : [0, \infty) \rightarrow \mathbb{R}^{n \times n}$ be a convolution-type kernel for a linear Volterra equation in $\mathbb{R}^n$. If $\forall \eta \in \mathbb{R}^n$, the function $t \mapsto \langle \eta, k(t)\eta \rangle$ is a nonnegative, nonincreasing and convex function on $(0, \infty)$, then k is nonnegative, meaning that $\forall \phi : [0, \infty) \rightarrow \mathbb{R}^n$ with compact support, there is*

$$\int_0^\infty \int_0^t \langle \phi(t), k(t-s)\phi(s) \rangle ds dt \geq 0.\tag{89}$$

Thus, to take advantage of this proposition, we need to verify that $\forall \eta \in \mathbb{R}^n$, $\langle \eta, \Gamma_t^\infty \eta \rangle$ is

(1) *nonnegative*:

$$\begin{aligned} & \langle \eta, \Gamma_t^\infty \eta \rangle \\ &= \int_{\Omega \times \Omega} \int_D \langle \nabla \varphi(\Theta_\infty(\theta), x), e^{-t \nabla \nabla V_\infty(\Theta_\infty(\theta))} \nabla \varphi(\Theta_\infty(\theta), x) \rangle \eta(x) \eta(x') \mu_0(d\theta) \hat{\nu}(dx) \hat{\nu}(dx') \\ &= \int_D \langle \mathbf{b}(\theta), e^{-t \nabla \nabla V_\infty(\Theta_\infty(\theta))} \mathbf{b}(\theta) \rangle \mu_0(d\theta) \geq 0, \end{aligned} \quad (90)$$

where

$$\mathbf{b}(\theta) = \int_{\Omega} \nabla \varphi(\Theta_\infty(\theta), x) \eta(x) \hat{\nu}(dx) \quad (91)$$

because by assumption, $\forall \theta \in D$, $\nabla \nabla V_\infty(\Theta_\infty(\theta))$ is positive semidefinite, and hence $e^{-t \nabla \nabla V_\infty(\Theta_\infty(\theta))}$ is a positive semidefinite operator;

(2) *nonincreasing*: Taking derivative with respect to time,

$$\frac{d}{dt} \langle \eta, \Gamma_t^\infty \eta \rangle = - \int_D \langle \mathbf{b}(\theta), \nabla \nabla V(\Theta_\infty(\theta)) e^{-t \nabla \nabla V_\infty(\Theta_\infty(\theta))} \mathbf{b}(\theta) \rangle \mu_0(d\theta) \leq 0, \quad (92)$$

because again, $\nabla \nabla V_\infty(\Theta_\infty(\theta))$ is positive semidefinite;

(3) *convex*: Taking one more derivative with respect to time,

$$\frac{d^2}{dt^2} \langle \eta, \Gamma_t^\infty \eta \rangle = \int_D \langle \mathbf{b}(\theta), (\nabla \nabla V(\Theta_\infty(\theta)))^2 e^{-t \nabla \nabla V_\infty(\Theta_\infty(\theta))} \mathbf{b}(\theta) \rangle \mu_0(d\theta) \geq 0, \quad (93)$$

Therefore, we can apply Proposition D.1 to conclude that Γ_{t-s}^∞ is PSD as a Volterra kernel, and so $\int_{t_0}^T \int_{t_0}^t \langle g_t, \Gamma_{t-s}^\infty g_s \rangle ds dt \geq 0$.

D.2 Proof of Theorem 3.5 (Unregularized case)

Recall that

$$\begin{aligned} \lim_{m \rightarrow \infty} m \mathbb{E}_0 \|f_t^{(m)} - f_t\|_{\hat{\nu}}^2 &= \mathbb{E}_0 \|g_t\|_{\hat{\nu}}^2 = \mathbb{E}_0 \int_{\Omega} \left| \int_D \varphi(\theta, x) \omega_t(d\theta) \right|^2 \hat{\nu}(dx) \\ &= \mathbb{E}_0 \int_{D \times D} K(\theta, \theta') \omega_t(d\theta) \omega_t(d\theta'), \end{aligned} \quad (94)$$

where, with a slight abuse of notation, in this equation $\mathbb{E}_0$ also denotes expectation over the randomness of the Gaussian distribution ω_0 defined in Proposition 3.1. From (31) in Proposition 3.1, this can be further expanded into

$$\begin{aligned} & \mathbb{E}_0 \int_{D \times D} K(\theta, \theta') \omega_t(d\theta) \omega_t(d\theta') \\ &= \mathbb{E}_0 \int_{D \times D} \langle T_t(\theta), \nabla \nabla' K(\Theta_t(\theta), \Theta_t(\theta')) T_t(\theta') \rangle \mu_0(d\theta) \mu_0(d\theta') \\ &+ 2 \mathbb{E}_0 \int_{D \times D} \nabla K(\Theta_t(\theta), \Theta_t(\theta')) T_t(\theta) \mu_0(d\theta) \omega_0(d\theta') \\ &+ \mathbb{E}_0 \int_{D \times D} K(\Theta_t(\theta), \Theta_t(\theta')) \omega_0(d\theta) \omega_0(d\theta'). \end{aligned} \quad (95)$$

The last term at the RHS is equal to $\mathbb{E}_0 \|\bar{g}_t\|_{\mathcal{V}}^2$ with $\bar{g}_t$ defined in (37). Using (32), it can be explicitly computed as

$$\begin{aligned}
\mathbb{E}_0 \|\bar{g}_t\|_{\mathcal{V}}^2 &= \mathbb{E}_0 \int_{D \times D} K(\Theta_t(\theta), \Theta_t(\theta')) \omega_0(d\theta) \omega_0(d\theta') \\
&= \int_{D \times D} K(\Theta_t(\theta), \Theta_t(\theta')) (\mu_0(d\theta) \delta_{\theta}(d\theta') - \mu_0(d\theta) \mu_0(d\theta')) \\
&= \int_D K(\theta, \theta) \mu_t(d\theta) - \int_{D \times D} K(\theta, \theta) \mu_t(d\theta) \mu_t(d\theta') \\
&= \int_D K(\theta, \theta) \mu_t(d\theta) - \|f_t\|_{\mathcal{V}}^2.
\end{aligned} \tag{96}$$

Thus,

$$\begin{aligned}
\lim_{t \rightarrow \infty} \mathbb{E}_0 \|\bar{g}_t\|_{\mathcal{V}}^2 &= \lim_{t \rightarrow \infty} \int_D K(\theta, \theta) \mu_t(d\theta) - \|f_t\|_{\mathcal{V}}^2 \\
&= \int_D K(\theta, \theta) \mu_{\infty}(d\theta) - \|f_{\infty}\|_{\mathcal{V}}^2 \\
&= \mathbb{E}_0 \|\bar{g}_{\infty}\|_{\mathcal{V}}^2
\end{aligned} \tag{97}$$

and so

$$\lim_{T \rightarrow \infty} \int_0^T \mathbb{E}_0 \|\bar{g}_t\|_{\mathcal{V}}^2 dt = \mathbb{E}_0 \|\bar{g}_{\infty}\|_{\mathcal{V}}^2, \tag{98}$$

where here and below we denote $f_0^t[\cdot] dt = \frac{1}{t} \int_0^t [\cdot] dt$. As a result, to prove (46) or (49) in Theorem 3.5, it suffices to establish that

$$\lim_{T \rightarrow \infty} \int_0^T \mathfrak{D}_t dt \leq 0, \tag{99}$$

or

$$\lim_{T \rightarrow \infty} \int_0^T \mathfrak{D}_t dt \leq -\mathbb{E}_0 \|\bar{g}_{\infty}\|_{\mathcal{V}}^2, \tag{100}$$

respectively, where we defined

$$\begin{aligned}
\mathfrak{D}_t &:= \mathbb{E}_0 \int_{D \times D} K(\theta, \theta') \omega_t(d\theta) \omega_t(d\theta') - \mathbb{E}_0 \int_{D \times D} K(\Theta_t(\theta), \Theta_t(\theta')) \omega_0(d\theta) \omega_0(d\theta') \\
&= \mathbb{E}_0 \int_{D \times D} \langle \mathbf{T}_t(\theta), \nabla \nabla' K(\Theta_t(\theta), \Theta_t(\theta')) \mathbf{T}_t(\theta') \rangle \mu_0(d\theta) \mu_0(d\theta') \\
&\quad + 2\mathbb{E}_0 \int_{D \times D} \nabla K(\Theta_t(\theta), \Theta_t(\theta')) \mathbf{T}_t(\theta) \mu_0(d\theta) \omega_0(d\theta').
\end{aligned} \tag{101}$$

To this end, we examine (33) as an infinite-dimensional ODE. With the Hilbert space $\mathcal{V}(D)$ defined in Appendix A and $\mathbf{b}_t$, $\mathcal{A}_t^{(K)}$ and $\mathcal{A}_t^{(V)}$ defined by (61), (63) and (64), respectively, we can rewrite (33) as the following ODE on $\mathcal{V}(D)$:

$$\dot{\mathbf{T}}_t = -(\mathcal{A}_t^{(K)} + \mathcal{A}_t^{(V)}) \mathbf{T}_t - \mathbf{b}_t, \tag{102}$$

We can also rewrite (101) as

$$\mathfrak{D}_t = \langle \mathbf{T}_t, \mathcal{A}_t^{(K)} \mathbf{T}_t \rangle_0 + 2\langle \mathbf{T}_t, \mathbf{b}_t \rangle_0. \tag{103}$$

From (102), we can deduce that

$$\frac{1}{2} \frac{d}{dt} \|\mathbf{T}_t\|_0^2 = -\langle \mathbf{T}_t, \mathcal{A}_t^{(V)} \mathbf{T}_t \rangle_0 - \langle \mathbf{T}_t, \mathcal{A}_t^{(K)} \mathbf{T}_t \rangle_0 - \langle \mathbf{T}_t, \mathbf{b}_t \rangle_0, \quad (104)$$

or equivalently

$$\langle \mathbf{T}_t, \mathcal{A}_t^{(K)} \mathbf{T}_t \rangle_0 + \langle \mathbf{T}_t, \mathbf{b}_t \rangle_0 = -\frac{1}{2} \frac{d}{dt} \|\mathbf{T}_t\|_0^2 - \langle \mathbf{T}_t, \mathcal{A}_t^{(V)} \mathbf{T}_t \rangle_0. \quad (105)$$

Therefore, we can rewrite (101) as

$$\begin{aligned} \mathfrak{D}_t &= 2 \left(\langle \mathbf{T}_t, \mathcal{A}_t^{(K)} \mathbf{T}_t \rangle_0 + \langle \mathbf{T}_t, \mathbf{b}_t \rangle_0 \right) - \langle \mathbf{T}_t, \mathcal{A}_t^{(K)} \mathbf{T}_t \rangle_0 \\ &= 2 \left(-\frac{1}{2} \frac{d}{dt} \|\mathbf{T}_t\|_0^2 - \langle \mathbf{T}_t, \mathcal{A}_t^{(V)} \mathbf{T}_t \rangle_0 \right) - \langle \mathbf{T}_t, \mathcal{A}_t^{(K)} \mathbf{T}_t \rangle_0 \\ &= -\frac{d}{dt} \|\mathbf{T}_t\|_0^2 - 2 \langle \mathbf{T}_t, \mathcal{A}_t^{(V)} \mathbf{T}_t \rangle_0 - \langle \mathbf{T}_t, \mathcal{A}_t^{(K)} \mathbf{T}_t \rangle_0 \end{aligned} \quad (106)$$

and as a result, since $\mathbf{T}_0 = 0$,

$$\int_0^T \mathfrak{D}_t dt = -\frac{1}{T} \|\mathbf{T}_T\|_0^2 - 2 \int_0^T \langle \mathbf{T}_t, \mathcal{A}_t^{(V)} \mathbf{T}_t \rangle_0 dt - \int_0^T \langle \mathbf{T}_t, \mathcal{A}_t^{(K)} \mathbf{T}_t \rangle_0 dt. \quad (107)$$

Note that for all t , $\mathcal{A}_t^{(K)}$ is a positive semidefinite (PSD) operator on $\mathcal{V}(D)$, as $\forall \xi \in \mathcal{V}(D)$,

$$\begin{aligned} \langle \mathcal{A}_t^{(K)} \xi, \xi \rangle_0 &= \mathbb{E}_0 \int_{D \times D} \langle \xi(\theta), \nabla \nabla' K(\Theta_t(\theta), \Theta_t(\theta')) \xi(\theta') \rangle_{\mu_0(d\theta) \mu_0(d\theta')} \\ &= \mathbb{E}_0 \int_{\Omega} \left| \int_D \nabla \varphi(\Theta_t(\theta)) \cdot \xi(\theta) \mu_0(d\theta) \right|^2 \hat{\nu}(d\mathbf{x}) \geq 0. \end{aligned} \quad (108)$$

This implies that $\int_0^T \langle \mathbf{T}_t, \mathcal{A}_t^{(K)} \mathbf{T}_t \rangle_0 dt \geq 0$. Hence, to establish (99), it is sufficient to show that

$$\lim_{T \rightarrow \infty} \int_0^T \langle \mathbf{T}_t, \mathcal{A}_t^{(V)} \mathbf{T}_t \rangle_0 dt = 0. \quad (109)$$

To this end, we need two lemmas that are proved below in Appendices D.2.1 and D.2.2, respectively:

Lemma D.2. *Assuming (47) and (48) together with Assumptions 2.2, 2.3 and 2.5, we have*

$$\int_0^\infty \|\mathcal{A}_t^{(V)}\|_0 dt < \infty \quad (110)$$

$$\int_0^\infty \|\mathcal{A}_\infty^{(K)} - \mathcal{A}_t^{(K)}\|_0 dt < \infty \quad (111)$$

$$\int_0^\infty \|\mathbf{b}_t - \mathbf{b}_\infty\|_0 dt < \infty \quad (112)$$

Lemma D.3. *Assuming (47) and (48) together with Assumptions 2.2, 2.3 and 2.5, we have*

$$\sup_{t < \infty} \|\mathbf{T}_t\|_0^2 < \infty. \quad (113)$$

With these two lemmas, we can show that

$$\begin{aligned} \left| \int_0^T \langle \mathbf{T}_t, \mathcal{A}_t^{(V)} \mathbf{T}_t \rangle_0 dt \right| &\leq \int_0^T \|\mathcal{A}_t^{(V)}\|_0 \|\mathbf{T}_t\|_0^2 dt \\ &\leq \left(\int_0^T \|\mathcal{A}_t^{(V)}\|_0 dt \right) \sup_{t < \infty} \|\mathbf{T}_t\|_0^2 \\ &< \infty, \end{aligned} \quad (114)$$

and therefore (109) is satisfied. This finishes the proof of (46) under (47) and (48) together with Assumptions 2.2, 2.3 and 2.5.

Next, we show (49) under the additional condition of Assumption 2.1. Thanks to (107) and (109), it is sufficient to establish that

$$\lim_{T \rightarrow \infty} \int_0^T \langle \mathbf{T}_t, \mathcal{A}_t^{(K)} \mathbf{T}_t \rangle_0 dt = \mathbb{E}_0 \|\bar{g}_\infty\|_\nu^2. \quad (115)$$

Heuristically, if $\mathbf{T}_\infty := \lim_{t \rightarrow \infty} \mathbf{T}_t$ exists, then from (102), it has to satisfy

$$-\mathbf{b}_\infty = \left(\mathcal{A}_\infty^{(V)} + \mathcal{A}_\infty^{(K)} \right) \mathbf{T}_\infty = \mathcal{A}_\infty^{(K)} \mathbf{T}_\infty, \quad (116)$$

as $\mathcal{A}_\infty^{(V)} = 0$ (because $\nabla \nabla V(\boldsymbol{\theta}, \mu_\infty) = \int_\Omega \varphi(\boldsymbol{\theta}, \mathbf{x})(f_\infty(\mathbf{x}) - f_*(\mathbf{x})) \hat{\nu}(d\mathbf{x}) = 0$ under the assumption of (47)). This equation implies that

$$(\mathbf{T}_\infty)^\parallel = - \left(\mathcal{A}_\infty^{(K)} \right)^\dagger \mathbf{b}_\infty, \quad (117)$$

where $(\mathbf{T}_\infty)^\parallel$ denotes the component of $\mathbf{T}_\infty$ in the range of $\mathcal{A}_\infty^{(K)}$, and $\left(\mathcal{A}_\infty^{(K)} \right)^\dagger$ denotes the Moore-Penrose pseudoinverse of $\mathcal{A}_\infty^{(K)}$. As a result,

$$\begin{aligned} \langle \mathbf{T}_\infty, \mathcal{A}_\infty^{(K)} \mathbf{T}_\infty \rangle_0 &= \langle (\mathbf{T}_\infty)^\parallel, \mathcal{A}_\infty^{(K)} (\mathbf{T}_\infty)^\parallel \rangle_0 \\ &= \langle - \left(\mathcal{A}_\infty^{(K)} \right)^\dagger \mathbf{b}_\infty, - \mathcal{A}_\infty^{(K)} \left(\mathcal{A}_\infty^{(K)} \right)^\dagger \mathbf{b}_\infty \rangle_0 \\ &= \langle \mathbf{b}_\infty, \left(\mathcal{A}_\infty^{(K)} \right)^\dagger \mathbf{b}_\infty \rangle_0. \end{aligned} \quad (118)$$

Rigorously, without assuming the existence of $\mathbf{T}_\infty$, we can establish that

Lemma D.4. *Assuming (47) and (48) together with Assumptions 2.2, 2.3 and 2.5, we have*

$$\lim_{t \rightarrow \infty} \int_0^t \langle \mathbf{T}_s, \mathcal{A}_s^{(K)} \mathbf{T}_s \rangle_0 ds \geq \langle \mathbf{b}_\infty, \left(\mathcal{A}_\infty^{(K)} \right)^\dagger \mathbf{b}_\infty \rangle_0. \quad (119)$$

As a consequence,

$$\lim_{t \rightarrow \infty} \int_0^t \mathbb{E}_0 \|g_s\|_\nu^2 dt \leq \mathbb{E}_0 \|\bar{g}_\infty\|_\nu^2 - \langle \mathbf{b}_\infty, \left(\mathcal{A}_\infty^{(K)} \right)^\dagger \mathbf{b}_\infty \rangle_0. \quad (120)$$

This lemma is proved in D.2.3. It implies that we only need to show that

$$\langle \mathbf{b}_\infty, \left(\mathcal{A}_\infty^{(K)} \right)^\dagger \mathbf{b}_\infty \rangle_0 = \mathbb{E}_0 \|\bar{g}_\infty\|_\nu^2. \quad (121)$$

This requires us to further exploit the relationship among $\mathcal{A}_\infty^{(K)}$, $\mathbf{b}_\infty$ and $\bar{g}_\infty$. With the Hilbert space $\mathcal{W}_L(\Omega)$ defined in Appendix A and $\mathcal{B}_t$ defined by (67), we can rewrite (63) as

$$\mathcal{A}_t^{(K)} = \mathcal{B}_t \mathcal{B}_t^\top. \quad (122)$$

Further, recall that

$$g_t = \int_D \varphi(\boldsymbol{\theta}, \cdot) \omega_t(d\boldsymbol{\theta}) = \int_D \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \cdot) \omega_0(d\boldsymbol{\theta}) + \int_D \nabla \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \cdot) \cdot \mathbf{T}_t(\boldsymbol{\theta}) \mu_0(d\boldsymbol{\theta}) \quad (123)$$

$$\bar{g}_t = \int_D \varphi(\boldsymbol{\theta}, \cdot) \bar{\omega}_t(d\boldsymbol{\theta}) = \int_D \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \cdot) \omega_0(d\boldsymbol{\theta}). \quad (124)$$

Therefore, we can write

$$g_t = \bar{g}_t + \mathcal{B}_t^\top \mathbf{T}_t, \quad (125)$$

and

$$\mathbf{b}_t = \mathcal{B}_t \bar{g}_t, \quad (126)$$

Similar formulas hold when we replace t by ∞ . With these relations, we see that

$$\begin{aligned} \left\langle \mathbf{b}_\infty, \left(\mathcal{A}_\infty^{(K)} \right)^\dagger \mathbf{b}_\infty \right\rangle_0 &= \left\langle \mathcal{B}_\infty \bar{g}_\infty, (\mathcal{B}_\infty \mathcal{B}_\infty^\top)^\dagger \mathcal{B}_\infty \bar{g}_\infty \right\rangle_0 \\ &= \mathbb{E}_0 \| (\mathcal{B}_\infty)^\dagger \mathcal{B}_\infty \bar{g}_\infty \|_\nu^2, \end{aligned} \quad (127)$$

because $(\mathcal{B}_\infty \mathcal{B}_\infty^\top)^\dagger = (\mathcal{B}_\infty^\top)^\dagger (\mathcal{B}_\infty)^\dagger$. Since $(\mathcal{B}_\infty)^\dagger \mathcal{B}_\infty$ is the projection operator (matrix) onto the range of $\mathcal{B}_\infty^\top$ in $\mathbb{R}^n$, it is then sufficient to prove that

Lemma D.5. *Under Assumptions 2.1, 2.2, 2.3 and 2.5, $\mathbb{P}_0$ -almost surely, $\bar{g}_\infty \in \text{Ran}(\mathcal{B}_\infty^\top)$.*

Lemma D.5 is proven in Appendix D.2.4 and it concludes the proof of (49) in Theorem 3.5.

To show that $\|g_t\|_\nu$ decreases monotonically when $\mu_0 = \mu_\infty$, note that in this case $\mu_t = \mu_\infty$, $\forall t \geq 0$, and so $\mathcal{A}_t^{(V)} = \mathcal{A}_\infty^{(V)} = 0$, $\mathcal{A}_t^{(K)} = \mathcal{A}_\infty^{(K)}$ and $\mathbf{b}_t = \mathbf{b}_\infty$, $\forall t \geq 0$. Thus, (102) becomes

$$\dot{\mathbf{T}}_t = -\mathcal{A}_\infty^{(K)} \mathbf{T}_t - \mathbf{b}_\infty, \quad (128)$$

As will be shown in Lemma D.6, $\mathbf{b}_\infty$ is in the range of $\mathcal{A}_\infty^{(K)}$. Therefore, defining

$$\mathbf{u}_\infty = (\mathcal{A}_\infty^{(K)})^\dagger \mathbf{b}_\infty, \quad (129)$$

and

$$\mathbf{z}_t = \mathbf{T}_t + \mathbf{u}_\infty, \quad (130)$$

there is

$$\dot{\mathbf{z}}_t = -\mathcal{A}_\infty^{(K)} \mathbf{z}_t, \quad (131)$$

whose solution can be written analytically as

$$\mathbf{z}_t = e^{-t\mathcal{A}_\infty^{(K)}} \mathbf{z}_0 = e^{-t\mathcal{A}_\infty^{(K)}} \mathbf{u}_\infty. \quad (132)$$

Thus,

$$\mathbf{T}_t = \mathbf{z}_t - \mathbf{u}_\infty = -(I - e^{-t\mathcal{A}_\infty^{(K)}}) \mathbf{u}_\infty \quad (133)$$

Therefore, as $\mathbf{b}_\infty = \mathcal{B}_\infty \bar{g}_\infty$, there is

$$\begin{aligned} g_t &= \bar{g}_\infty + \mathcal{B}_\infty^\top \mathbf{T}_t \\ &= \bar{g}_\infty - \mathcal{B}_\infty^\top (I - e^{-t\mathcal{A}_\infty^{(K)}}) \mathbf{u}_\infty \\ &= \bar{g}_\infty - \mathcal{B}_\infty^\top (I - e^{-t\mathcal{A}_\infty^{(K)}}) (\mathcal{A}_\infty^{(K)})^\dagger \mathcal{B}_\infty \bar{g}_\infty . \end{aligned} \quad (134)$$

Hence,

$$|g_\infty|^2 = |\bar{g}_\infty|^2 - 2(*) + (**), \quad (135)$$

where

$$\begin{aligned} (*) &= (\mathcal{B}_\infty \bar{g}_\infty)^\top (I - e^{-t\mathcal{A}_\infty^{(K)}}) (\mathcal{A}_\infty^{(K)})^\dagger \mathcal{B}_\infty \bar{g}_\infty \\ &= \mathbf{b}_\infty^\top (I - e^{-t\mathcal{A}_\infty^{(K)}}) (\mathcal{A}_\infty^{(K)})^\dagger \mathbf{b}_\infty \end{aligned} \quad (136)$$

and

$$\begin{aligned} (**) &= (\mathcal{B}_\infty \bar{g}_\infty)^\top (\mathcal{A}_\infty^{(K)})^\dagger (I - e^{-t\mathcal{A}_\infty^{(K)}}) \mathcal{B}_\infty \mathcal{B}_\infty^\top (I - e^{-t\mathcal{A}_\infty^{(K)}}) (\mathcal{A}_\infty^{(K)})^\dagger \mathcal{B}_\infty \bar{g}_\infty \\ &= \mathbf{b}_\infty^\top (I - e^{-t\mathcal{A}_\infty^{(K)}}) \mathcal{B}_\infty \mathcal{B}_\infty^\top (I - e^{-t\mathcal{A}_\infty^{(K)}}) \mathbf{b}_\infty . \end{aligned} \quad (137)$$

In the ERM setting, $\mathcal{A}_\infty^{(K)}$ is PSD with a finite number of nonzero eigenspaces. Consider a set of its orthonormal eigenfunctions that span those nonzero eigenspaces, $v_1, \dots, v_k$, corresponding to eigenvalues $\lambda_1, \dots, \lambda_k > 0$, respectively. As $\mathbf{b}_\infty$ is in the range of $\mathcal{A}_\infty^{(K)}$ by Lemma D.6, we can decompose it as

$$\mathbf{b}_\infty = \sum_{i=1}^k c_i v_i \quad (138)$$

for some real numbers c_i 's. Thus, we can write

$$\begin{aligned} (*) &= \left(\sum_{i=1}^k c_i v_i \right)^\top (I - e^{-t\mathcal{A}_\infty^{(K)}}) (\mathcal{A}_\infty^{(K)})^\dagger \left(\sum_{j=1}^k c_j v_j \right) \\ &= \left(\sum_{i=1}^k c_i v_i \right)^\top \left(\sum_{j=1}^k c_j \lambda_j^{-1} (1 - e^{-\lambda_j t}) v_j \right) \\ &= \sum_{i=1}^k \lambda_i^{-1} (1 - e^{-\lambda_i t}) c_i^2 , \end{aligned} \quad (139)$$

$$\begin{aligned} (**) &= \left(\sum_{i=1}^k c_i v_i \right)^\top (\mathcal{A}_\infty^{(K)})^\dagger (I - e^{-t\mathcal{A}_\infty^{(K)}}) \mathcal{B}_\infty \mathcal{B}_\infty^\top (I - e^{-t\mathcal{A}_\infty^{(K)}}) (\mathcal{A}_\infty^{(K)})^\dagger \left(\sum_{j=1}^k c_j v_j \right) \\ &= \left(\sum_{i=1}^k c_i v_i \right)^\top (\mathcal{A}_\infty^{(K)})^\dagger (I - e^{-t\mathcal{A}_\infty^{(K)}}) \mathcal{A}_\infty^{(K)} (I - e^{-t\mathcal{A}_\infty^{(K)}}) (\mathcal{A}_\infty^{(K)})^\dagger \left(\sum_{j=1}^k c_j v_j \right) \\ &= \left(\sum_{i=1}^k c_i v_i \right)^\top \left(\sum_{j=1}^k \lambda_j^{-1} (1 - e^{-\lambda_j t})^2 c_j v_j \right) \\ &= \sum_{i=1}^k \lambda_i^{-1} (1 - e^{-\lambda_i t})^2 c_i^2 . \end{aligned} \quad (140)$$

Therefore,

$$\begin{aligned}
|g_\infty|^2 &= |\bar{g}_\infty|^2 - 2 \sum_{i=1}^k \lambda_j^{-1} (1 - e^{-\lambda_j t}) c_i^2 + \sum_{i=1}^k \lambda_j^{-1} (1 - e^{-\lambda_j t})^2 c_i^2 \\
&= |\bar{g}_\infty|^2 + \sum_{i=1}^k \lambda_j^{-1} (1 - e^{-\lambda_j t}) (-1 - e^{-\lambda_j t}) c_i^2 \\
&= |\bar{g}_\infty|^2 - \sum_{i=1}^k \lambda_j^{-1} (1 - e^{-2\lambda_j t}) c_i^2,
\end{aligned} \tag{141}$$

which is decreasing in time. This completes the proof of Theorem 3.5. $\square$

D.2.1 Proof of Lemma D.2

Proof of (110): $\int_0^\infty \|\mathcal{A}_t^{(V)}\|_0 dt < \infty$.

By the definition of the operator norm induced by $\|\cdot\|_0$ on $\mathcal{V}(D)$, $\|\mathcal{A}_t^{(V)}\|_0$ is the smallest number C_t such that $\forall \xi$, there is

$$\|\mathcal{A}_t^{(V)}\|_0 = \sup_{\xi \in \mathcal{V}(D), \|\xi\|_0 \neq 0} \frac{|\langle \xi, \mathcal{A}_t^{(V)} \xi \rangle_0|}{\|\xi\|_0^2}. \tag{142}$$

In the unregularized case, a straightforward bound of $|\langle \xi, \mathcal{A}_t^{(V)} \xi \rangle_0|$ is

$$\begin{aligned}
|\langle \xi, \mathcal{A}_t^{(V)} \xi \rangle_0| &= \left| \mathbb{E}_0 \int_D \langle \xi(\theta), \nabla \nabla V(\Theta_t(\theta), \mu_t) \xi(\theta) \rangle \mu_0(d\theta) \right| \\
&= \left| \mathbb{E}_0 \int_D \int_\Omega \langle \xi(\theta), \nabla \nabla \varphi(\Theta_t(\theta), x) \xi(\theta) \rangle (f_t(x) - f_*(x)) \hat{\nu}(dx) \mu_0(d\theta) \right| \\
&\leq \mathbb{E}_0 \int_D \int_\Omega C_{\nabla \nabla \varphi} |\xi(\theta)|^2 |f_t(x) - f_*(x)| \hat{\nu}(dx) \mu_0(d\theta) \\
&= C_{\nabla \nabla \varphi} \|\xi\|_0^2 \int_\Omega |f_t(x) - f_*(x)| \hat{\nu}(dx) \\
&\leq n^{1/2} C_{\nabla \nabla \varphi} \|\xi\|_0^2 \|f_t - f_*\|_{\hat{\nu}} \\
&= n^{1/2} C_{\nabla \nabla \varphi} \|\xi\|_0^2 (\mathcal{L}(\mu_t))^{1/2}.
\end{aligned} \tag{143}$$

Thus, we have

$$\|\mathcal{A}_t^{(V)}\|_0 \leq n^{1/2} C_{\nabla \nabla \varphi} \|\xi\|_0^2 (\mathcal{L}(\mu_t))^{1/2}. \tag{144}$$

By the assumption (48), we thus have

$$\int_0^\infty \|\mathcal{A}_t^{(V)}\|_0 dt \leq n^{1/2} C_{\nabla \nabla \varphi} \int_0^\infty (\mathcal{L}(\mu_t))^{1/2} dt < \infty \tag{145}$$

which gives us the desired bound. $\square$

Proof of (111): $\int_0^\infty \|\mathcal{A}_\infty^{(K)} - \mathcal{A}_t^{(K)}\|_0 dt < \infty$.

We have

$$\begin{aligned}
& \langle \xi, (\mathcal{A}_t^{(K)} - \mathcal{A}_\infty^{(K)}) \xi \rangle_0 \\
&= \mathbb{E}_0 \int_{\Omega} \left(\left(\int_D \nabla \varphi(\Theta_t(\theta), x) \cdot \xi(\theta) \mu_0(d\theta) \right)^2 - \left(\int_D \nabla \varphi(\Theta_\infty(\theta), x) \cdot \xi(\theta) \mu_0(d\theta) \right)^2 \right) \hat{\nu}(dx) \\
&= \mathbb{E}_0 \int_{\Omega} \left(\int_D (\nabla \varphi(\Theta_t(\theta), x) + \nabla \varphi(\Theta_\infty(\theta), x)) \cdot \xi(\theta) \mu_0(d\theta) \right) \\
&\quad \times \left(\int_D (\nabla \varphi(\Theta_t(\theta), x) - \nabla \varphi(\Theta_\infty(\theta), x)) \cdot \xi(\theta) \mu_0(d\theta) \right) \hat{\nu}(dx). \tag{146}
\end{aligned}$$

Hence, the absolute value of the expression above is upper-bounded by

$$\begin{aligned}
& \mathbb{E}_0 \left(\int_D |\nabla \varphi(\Theta_t(\theta), x) + \nabla \varphi(\Theta_\infty(\theta), x)| |\xi(\theta)| \mu_0(d\theta) \right. \\
&\quad \times \left. \int_D |\nabla \varphi(\Theta_t(\theta), x) - \nabla \varphi(\Theta_\infty(\theta), x)| |\xi(\theta)| \mu_0(d\theta) \right) \\
&\leq 2C_{\nabla \varphi} C_{\nabla \nabla \varphi} \|\xi\|_0^2 \left(\int_D |\Theta_t(\theta) - \Theta_\infty(\theta)|^2 \mu_0(d\theta) \right)^{1/2}. \tag{147}
\end{aligned}$$

Thus, by the assumption (48), we have

$$\begin{aligned}
\int_0^\infty \|\mathcal{A}_\infty^{(K)} - \mathcal{A}_t^{(K)}\|_0 dt &\leq 2C_{\nabla \varphi} C_{\nabla \nabla \varphi} \int_0^\infty \left(\int_D |\Theta_t(\theta) - \Theta_\infty(\theta)|^2 \mu_0(d\theta) \right)^{1/2} dt \\
&\leq 2C_{\nabla \varphi} C_{\nabla \nabla \varphi} \int_0^\infty (\mathcal{L}(\mu_t))^{1/2} dt \\
&< \infty \tag{148}
\end{aligned}$$

□

Proof of (112): $\int_0^\infty \|\mathbf{b}_t - \mathbf{b}_\infty\|_0 dt < \infty$.

There is

$$\begin{aligned}
& \mathbf{b}_t(\theta) - \mathbf{b}_\infty(\theta) \\
&= \int_D (\nabla K(\Theta_t(\theta), \Theta_t(\theta')) - \nabla K(\Theta_\infty(\theta), \Theta_\infty(\theta'))) \omega_0(d\theta') \\
&= \int_D \int_{\Omega} \nabla \varphi(\Theta_t(\theta), x) \cdot \nabla \varphi(\Theta_t(\theta'), x)^\top - \nabla \varphi(\Theta_\infty(\theta), x) \cdot \nabla \varphi(\Theta_\infty(\theta'), x)^\top \hat{\nu}(dx) \omega_0(d\theta') \\
&= \int_D \int_{\Omega} \nabla \varphi(\Theta_t(\theta), x) \cdot \nabla \varphi(\Theta_t(\theta'), x)^\top - \nabla \varphi(\Theta_t(\theta), x) \cdot \nabla \varphi(\Theta_\infty(\theta'), x)^\top \hat{\nu}(dx) \omega_0(d\theta') \\
&\quad + \int_D \int_{\Omega} \nabla \varphi(\Theta_t(\theta), x) \cdot \nabla \varphi(\Theta_\infty(\theta'), x)^\top - \nabla \varphi(\Theta_\infty(\theta), x) \cdot \nabla \varphi(\Theta_\infty(\theta'), x)^\top \omega_0(d\theta') \\
&= \int_{\Omega} \nabla \varphi(\Theta_t(\theta), x) \cdot \left(\int_D (\nabla \varphi(\Theta_t(\theta'), x) - \nabla \varphi(\Theta_\infty(\theta'), x)) \omega_0(d\theta') \right)^\top \hat{\nu}(dx) \\
&\quad + \int_{\Omega} (\nabla \varphi(\Theta_t(\theta), x) - \nabla \varphi(\Theta_\infty(\theta), x)) \left(\int_D \nabla \varphi(\Theta_\infty(\theta'), x) \omega_0(d\theta') \right)^\top \hat{\nu}(dx). \tag{149}
\end{aligned}$$

Thus,

$$\begin{aligned}
& \mathbb{E}_0 |\mathbf{b}_t(\boldsymbol{\theta}) - \mathbf{b}_\infty(\boldsymbol{\theta})|^2 \\
& \leq \mathbb{E}_0 \left| \int_{\Omega} \nabla \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mathbf{x}) \cdot \left(\int_D (\nabla \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}'), \mathbf{x}) - \nabla \varphi(\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}'), \mathbf{x})) \omega_0(d\boldsymbol{\theta}') \right)^\top \hat{\nu}(d\mathbf{x}) \right|^2 \\
& \quad + \mathbb{E}_0 \left| \int_{\Omega} (\nabla \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mathbf{x}) - \nabla \varphi(\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}), \mathbf{x})) \left(\int_D \nabla \varphi(\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}'), \mathbf{x}) \omega_0(d\boldsymbol{\theta}') \right)^\top \hat{\nu}(d\mathbf{x}) \right|^2 \\
& \leq \int_{\Omega} |\nabla \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mathbf{x})|^2 \mathbb{E}_0 \left| \int_D (\nabla \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}'), \mathbf{x}) - \nabla \varphi(\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}'), \mathbf{x})) \omega_0(d\boldsymbol{\theta}') \right|^2 \hat{\nu}(d\mathbf{x}) \\
& \quad + \int_{\Omega} |\nabla \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mathbf{x}) - \nabla \varphi(\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}), \mathbf{x})|^2 \mathbb{E}_0 \left| \int_D \nabla \varphi(\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}'), \mathbf{x}) \omega_0(d\boldsymbol{\theta}') \right|^2 \hat{\nu}(d\mathbf{x}) \\
& \leq C_{\nabla \varphi}^2 \int_{\Omega} \mathbb{E}_0 \left| \int_D (\nabla \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}'), \mathbf{x}) - \nabla \varphi(\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}'), \mathbf{x})) \omega_0(d\boldsymbol{\theta}') \right|^2 \hat{\nu}(d\mathbf{x}) \\
& \quad + C_{\nabla \varphi}^2 |\boldsymbol{\Theta}_t(\boldsymbol{\theta}) - \boldsymbol{\Theta}_\infty(\boldsymbol{\theta})|^2 \int_{\Omega} \mathbb{E}_0 \left| \int_D \nabla \varphi(\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}'), \mathbf{x}) \omega_0(d\boldsymbol{\theta}') \right|^2 \hat{\nu}(d\mathbf{x}) .
\end{aligned} \tag{150}$$

By the property of ω_0 , there is

$$\begin{aligned}
\mathbb{E}_0 \left| \int_D \chi(\boldsymbol{\theta}) \omega_0(d\boldsymbol{\theta}) \right|^2 &= \int_D \left| \chi(\boldsymbol{\theta}) - \int_D \chi(\boldsymbol{\theta}') \mu_0(d\boldsymbol{\theta}') \right|^2 \mu_0(d\boldsymbol{\theta}) \\
&\leq \int_D |\chi(\boldsymbol{\theta})|^2 \mu_0(d\boldsymbol{\theta})
\end{aligned} \tag{151}$$

for a test function χ on D . Thus,

$$\begin{aligned}
\mathbb{E}_0 |\mathbf{b}_t(\boldsymbol{\theta}) - \mathbf{b}_\infty(\boldsymbol{\theta})|^2 &\leq C_{\nabla \varphi}^2 \int_{\Omega} \int_D |\nabla \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}'), \mathbf{x}) - \nabla \varphi(\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}'), \mathbf{x})|^2 \mu_0(d\boldsymbol{\theta}') \\
&\quad + C_{\nabla \varphi}^2 |\boldsymbol{\Theta}_t(\boldsymbol{\theta}) - \boldsymbol{\Theta}_\infty(\boldsymbol{\theta})|^2 \int_{\Omega} \int_D |\nabla \varphi(\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}'), \mathbf{x})|^2 \mu_0(d\boldsymbol{\theta}') \hat{\nu}(d\mathbf{x}) \\
&\leq C_{\nabla \varphi}^2 C_{\nabla \varphi}^2 \int_D |\boldsymbol{\Theta}_t(\boldsymbol{\theta}') - \boldsymbol{\Theta}_\infty(\boldsymbol{\theta}')|^2 \mu_0(d\boldsymbol{\theta}') \\
&\quad + C_{\nabla \varphi}^2 C_{\nabla \varphi}^2 |\boldsymbol{\Theta}_t(\boldsymbol{\theta}) - \boldsymbol{\Theta}_\infty(\boldsymbol{\theta})|^2 .
\end{aligned} \tag{152}$$

Therefore,

$$\begin{aligned}
\|\mathbf{b}_t - \mathbf{b}_\infty\|_0^2 &= \mathbb{E}_0 \int_D |\mathbf{b}_t(\boldsymbol{\theta}) - \mathbf{b}_\infty(\boldsymbol{\theta})|^2 \mu_0(d\boldsymbol{\theta}) \\
&\leq 2C_{\nabla \varphi}^2 C_{\nabla \varphi}^2 \int_D |\boldsymbol{\Theta}_t(\boldsymbol{\theta}) - \boldsymbol{\Theta}_\infty(\boldsymbol{\theta})|^2 \mu_0(d\boldsymbol{\theta}) .
\end{aligned} \tag{153}$$

Since

$$\begin{aligned}
\int_D |\boldsymbol{\Theta}_t(\boldsymbol{\theta}) - \boldsymbol{\Theta}_\infty(\boldsymbol{\theta})|^2 \mu_0(d\boldsymbol{\theta}) &\leq \int_D \int_t^\infty \left| \dot{\boldsymbol{\Theta}}_s(\boldsymbol{\theta}) \right|^2 ds \mu_0(d\boldsymbol{\theta}) \\
&= \int_D \int_t^\infty |\nabla V(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mu_t)|^2 ds \mu_0(d\boldsymbol{\theta}) \\
&= - \int_t^\infty \frac{d}{ds} \mathcal{L}(\mu_s) ds \\
&= \mathcal{L}(\mu_t) - \mathcal{L}(\mu_\infty) \\
&= \mathcal{L}(\mu_t)
\end{aligned} \tag{154}$$

we can conclude that

$$\int_0^\infty \|\mathbf{b}_t - \mathbf{b}_\infty\|_0 dt \leq \int_0^\infty |\mathcal{L}(\mu_t)|^{1/2} dt < \infty. \quad (155)$$

□

D.2.2 Proof of Lemma D.3

Our goal is to show that $\|\mathbf{T}_t\|_0$ remains bounded for all time. First note that, for all t , $\mathcal{A}_t^{(K)}$ is a positive semidefinite (PSD) operator on $\mathcal{V}(D)$ since

$$\begin{aligned} \langle \mathcal{A}_t^{(K)} \boldsymbol{\xi}, \boldsymbol{\xi} \rangle_0 &= \mathbb{E}_0 \int_{D \times D} \langle \boldsymbol{\xi}(\boldsymbol{\theta}), \nabla \nabla' K(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \boldsymbol{\Theta}_t(\boldsymbol{\theta}')) \boldsymbol{\xi}(\boldsymbol{\theta}') \rangle \mu_0(d\boldsymbol{\theta}) \mu_0(d\boldsymbol{\theta}') \\ &= \mathbb{E}_0 \int_\Omega \left| \int_D \nabla \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta})) \cdot \boldsymbol{\xi}(\boldsymbol{\theta}) \mu_0(d\boldsymbol{\theta}) \right|^2 \hat{\nu}(d\mathbf{x}) \geq 0. \end{aligned} \quad (156)$$

Second, by Assumption 2.5, for μ_0 -almost-every $\boldsymbol{\theta} \in D$, $\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}) = \lim_{t \rightarrow \infty} \boldsymbol{\Theta}_t(\boldsymbol{\theta})$ exists, which allows us to define $\mathbf{b}_\infty$, $\mathcal{A}_\infty^{(K)}$, and $\mathcal{A}_\infty^{(V)}$ similarly to (61), (63) and (64) by replacing $\boldsymbol{\Theta}_t(\cdot)$ with $\boldsymbol{\Theta}_\infty(\cdot)$. Since we assume that

$$\forall \mathbf{x}_k \in \text{supp } \hat{\nu} \quad : \quad f_\infty(\mathbf{x}_k) = \int_D \varphi(\boldsymbol{\theta}, \mathbf{x}_k) \mu_\infty(d\boldsymbol{\theta}) = f_*(\mathbf{x}_k) \quad (157)$$

we have

$$\forall \boldsymbol{\theta} \in D \quad : \quad \nabla \nabla V(\boldsymbol{\theta}, \mu_\infty) = \int_\Omega \nabla \nabla \varphi(\boldsymbol{\theta}, \mathbf{x}) (f_\infty(\mathbf{x}) - f_*(\mathbf{x})) d\mathbf{x} = 0. \quad (158)$$

This implies that $\mathcal{A}_\infty^{(V)}$ is the zero operator on $\mathcal{V}(D)$.

Third, we have the following observation:

Lemma D.6. *Under Assumptions 2.2, 2.3 and 2.5, $\mathbf{b}_t \in \text{Ran}(\mathcal{A}_t^{(K)})$ for all t , and $\mathbf{b}_\infty \in \text{Ran}(\mathcal{A}_\infty^{(K)})$. Specifically, $\exists \tilde{\mathbf{u}}_\infty \in \mathcal{V}(D)$ such that $\|\mathbf{u}_\infty\|_0 < \infty$ and $\mathcal{A}_\infty^{(K)} \tilde{\mathbf{u}}_\infty = \mathbf{b}_\infty$.*

Proof of Lemma D.6: Recall from (126) that $\mathbf{b}_\infty = \mathcal{B}_\infty \bar{g}_\infty$. Define $\tilde{\mathbf{u}}_\infty = \mathcal{B}_\infty (\mathcal{B}_\infty^\top \mathcal{B}_\infty)^\dagger \bar{g}_\infty$. We claim that $\mathcal{A}_\infty^{(K)} \tilde{\mathbf{u}}_\infty = \mathbf{b}_\infty$, because

$$\begin{aligned} \mathcal{A}_\infty^{(K)} \tilde{\mathbf{u}}_\infty &= (\mathcal{B}_\infty \mathcal{B}_\infty^\top) \mathcal{B}_\infty (\mathcal{B}_\infty^\top \mathcal{B}_\infty)^\dagger \bar{g}_\infty \\ &= \mathcal{B}_\infty \mathcal{B}_\infty^\top \left(\mathcal{B}_\infty (\mathcal{B}_\infty)^\dagger \right) (\mathcal{B}_\infty^\top)^\dagger \bar{g}_\infty \\ &= \mathcal{B}_\infty \left(\mathcal{B}_\infty^\top (\mathcal{B}_\infty^\top)^\dagger \right) \bar{g}_\infty \\ &= \mathcal{B}_\infty \bar{g}_\infty \\ &= \mathbf{b}_\infty, \end{aligned} \quad (159)$$

where the third equality is because $\mathcal{B}_\infty (\mathcal{B}_\infty)^\dagger$ is the projection operator onto $\text{Ran}(\mathcal{B}_\infty) = \text{Nul}^\perp(\mathcal{B}_\infty^\top)$, and the fourth equality is because $\mathcal{B}_\infty^\top (\mathcal{B}_\infty^\top)^\dagger$ is the projection operator onto $\text{Ran}(\mathcal{B}_\infty^\top) = \text{Nul}^\perp(\mathcal{B}_\infty)$.

It remains to establish that $\|\tilde{\mathbf{u}}_\infty\|_0 < \infty$. To show this, we see that

$$\begin{aligned}
& \int_D |\tilde{\mathbf{u}}_\infty(\boldsymbol{\theta})|^2 \mu_0(d\boldsymbol{\theta}) \\
&= \int_D \int_{\Omega \times \Omega} \left(\nabla \varphi(\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}), \mathbf{x}) (\mathcal{M}_\infty^\dagger \bar{g}_\infty)(\mathbf{x}) \right) \\
&\quad \cdot \left(\nabla \varphi(\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}), \mathbf{x}') ((\mathcal{M}_\infty)_\infty^\dagger \bar{g}_\infty)(\mathbf{x}') \right) \hat{\nu}(d\mathbf{x}) \hat{\nu}(d\mathbf{x}') \mu_0(d\boldsymbol{\theta}) \\
&= \int_\Omega \int_\Omega M(\mathbf{x}, \mathbf{x}', \mu_\infty) (\mathcal{M}_\infty^\dagger \bar{g}_\infty)(\mathbf{x}) (\mathcal{M}_\infty^\dagger \bar{g}_\infty)(\mathbf{x}') \hat{\nu}(d\mathbf{x}) \hat{\nu}(d\mathbf{x}') \\
&= \int_\Omega (\mathcal{M}_\infty^\dagger \bar{g}_\infty)(\mathbf{x}) \cdot \bar{g}_\infty(\mathbf{x}) \hat{\nu}(d\mathbf{x}) \\
&\leq \lambda_{\min}^{-1} \int_\Omega |\bar{g}_\infty(\mathbf{x})|^2 \hat{\nu}(d\mathbf{x}) ,
\end{aligned} \tag{160}$$

where $\lambda_{\min}$ is the least nonzero eigenvalue of the matrix $\mathcal{M}_\infty$ (and hence $\lambda_{\min}^{-1}$ is the largest eigenvalue of $\mathcal{M}_\infty^\dagger$). Since

$$\begin{aligned}
\mathbb{E}_0 |\bar{g}_\infty(\mathbf{x})|^2 &= \mathbb{E}_0 \left| \int_D \varphi(\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}), \mathbf{x}) \omega_0(d\boldsymbol{\theta}) \right|^2 \\
&= \int_D \left(\varphi(\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}), \mathbf{x}) - \int_D \varphi(\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}'), \mathbf{x}) \mu_0(d\boldsymbol{\theta}') \right)^2 \mu_0(d\boldsymbol{\theta}) \\
&\leq \int_D |\varphi(\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}), \mathbf{x})|^2 \mu_0(d\boldsymbol{\theta}) ,
\end{aligned} \tag{161}$$

there is

$$\begin{aligned}
\|\tilde{\mathbf{u}}_\infty\|_0^2 &\leq \mathbb{E}_0 \int_D |\tilde{\mathbf{u}}_\infty(\boldsymbol{\theta})|^2 \mu_0(d\boldsymbol{\theta}) \\
&\leq \lambda_{\min}^{-1} \int_\Omega \int_D (\varphi(\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}), \mathbf{x}))^2 \mu_0(d\boldsymbol{\theta}) \nu(d\mathbf{x}) \\
&\leq \lambda_{\min}^{-1} C_\varphi^2 < \infty ,
\end{aligned} \tag{162}$$

(End of the proof of Lemma D.6) $\square$

Coming back to the prof of Lemma D.3, we have shown that, as $t \rightarrow \infty$, (102) approaches the asymptotic dynamics

$$\dot{\mathbf{T}}_t = -\mathcal{A}_\infty^{(K)} \mathbf{T}_t - \mathbf{b}_\infty, \tag{163}$$

with $\mathcal{A}_\infty^{(K)}$ positive semidefinite and $\mathbf{b}_\infty$ in the range of $\mathcal{A}_\infty^{(K)}$. This is a stable system. Hence, the rest of the task is to examine what happens at finite time. To do so, we perform a change-of-variable with

$$\mathbf{z}_t = \mathbf{T}_t + \tilde{\mathbf{u}}_\infty, \tag{164}$$

with

$$\mathbf{u}_\infty = \mathcal{B}_\infty (\mathcal{B}_\infty^\top \mathcal{B}_\infty)^\dagger \bar{g}_\infty \tag{165}$$

as is defined in the proof of Lemma D.6. The dynamics of $\mathbf{z}_t$ is governed by

$$\begin{aligned}
\dot{\mathbf{z}}_t &= \dot{\mathbf{T}}_t = -(\mathcal{A}_t^{(K)} + \mathcal{A}_t^{(V)}) \mathbf{T}_t - \mathbf{b}_t \\
&= -\mathcal{A}_t^{(K)} \mathbf{z}_t - \mathcal{A}_t^{(V)} \mathbf{z}_t - (\mathbf{b}_t - (\mathcal{A}_t^{(K)} + \mathcal{A}_t^{(V)}) \tilde{\mathbf{u}}_\infty) .
\end{aligned} \tag{166}$$

Thus, in integral form,

$$\mathbf{z}_t = \Pi(t, 0)\mathbf{z}_0 + \int_0^t \Pi(t, s) \left(-\mathcal{A}_s^{(V)}\mathbf{z}_s - (\mathbf{b}_s - (\mathcal{A}_s^{(K)} + \mathcal{A}_s^{(V)})\tilde{\mathbf{u}}_\infty) \right) ds, \quad (167)$$

where $\Pi(t, s)$ is the fundamental solution (a.k.a. Green's function) associated with the time-variant homogeneous system

$$\dot{\mathbf{z}}_t = -\mathcal{A}_t^{(K)}\mathbf{z}_t. \quad (168)$$

Since $\mathcal{A}_t^{(K)}$ is positive semidefinite for all t , there is $\|\Pi(t, s)\|_0 \leq 1$ for $t > s$, where with a slight abuse of notation we also use $\|\cdot\|_0$ for the operator norm. Hence,

$$\begin{aligned} \|\mathbf{z}_t\|_0 &\leq \|\Pi(t, 0)\|_0 \|\mathbf{z}_0\|_0 + \int_0^t \|\Pi(t, s)\|_0 \left(\|\mathcal{A}_s^{(V)}\|_0 \|\mathbf{z}_s\|_0 + \|\mathbf{b}_s - (\mathcal{A}_s^{(K)} + \mathcal{A}_s^{(V)})\tilde{\mathbf{u}}_\infty\|_0 \right) ds \\ &\leq \|\mathbf{z}_0\|_0 + \int_0^t \left(\|\mathcal{A}_s^{(V)}\|_0 \|\mathbf{z}_s\|_0 + \|\mathbf{b}_s - (\mathcal{A}_s^{(K)} + \mathcal{A}_s^{(V)})\tilde{\mathbf{u}}_\infty\|_0 \right) ds. \end{aligned} \quad (169)$$

By Grönwall's inequality, we thus have

$$\|\mathbf{z}_t\|_0 \leq \left(\|\mathbf{z}_0\|_0 + \int_0^t \|\mathbf{b}_s - (\mathcal{A}_s^{(K)} + \mathcal{A}_s^{(V)})\tilde{\mathbf{u}}_\infty\|_0 ds \right) e^{\int_0^t \|\mathcal{A}_s^{(V)}\|_0 ds}. \quad (170)$$

Therefore, $\|\mathbf{z}_t\|_0$ remains bounded for all time if we can show that

$$\int_0^\infty \|\mathbf{b}_t - (\mathcal{A}_t^{(K)} + \mathcal{A}_t^{(V)})\tilde{\mathbf{u}}_\infty\|_0 dt < \infty, \quad \int_0^\infty \|\mathcal{A}_t^{(V)}\|_0 dt < \infty. \quad (171)$$

Since

$$\|\mathbf{b}_t - (\mathcal{A}_t^{(K)} + \mathcal{A}_t^{(V)})\tilde{\mathbf{u}}_\infty\|_0 \leq \|\mathbf{b}_t - \mathbf{b}_\infty\|_0 + \|(\mathcal{A}_t^{(K)} - \mathcal{A}_\infty^{(K)})\tilde{\mathbf{u}}_\infty\|_0 + \|\mathcal{A}_\infty^{(V)}\tilde{\mathbf{u}}_\infty\|_0 \quad (172)$$

we see that (171) is guaranteed by Lemmas D.2 and D.6.

This completes the proof of Lemma D.3. $\square$

D.2.3 Proof of Lemma D.4

From D.3, we have that

$$\lim_{t \rightarrow \infty} \left\| \int_0^t \dot{\mathbf{T}}_s ds \right\|_0 = \lim_{t \rightarrow \infty} \left\| \frac{1}{t} (\mathbf{T}_t - \mathbf{T}_0) \right\|_0 = 0. \quad (173)$$

By (102), we then obtain that

$$\lim_{t \rightarrow \infty} \left\| \int_0^t \left(\mathcal{A}_s^{(K)} \mathbf{T}_s + \mathbf{b}_s \right) ds + \int_0^t \mathcal{A}_s^{(V)} \mathbf{T}_s ds \right\|_0 = 0. \quad (174)$$

By (110) in Lemma D.2 as well as Lemma D.3, we know that

$$\lim_{t \rightarrow \infty} \left\| \int_0^t \mathcal{A}_s^{(V)} \mathbf{T}_s ds \right\|_0 = 0. \quad (175)$$

Therefore,

$$\lim_{t \rightarrow \infty} \left\| \int_0^t \left(\mathcal{A}_s^{(K)} \mathbf{T}_s + \mathbf{b}_s \right) ds \right\|_0 = 0. \quad (176)$$

Next, by (111) and (112) in Lemma D.2 as well as Lemma D.3, we know that

$$\lim_{t \rightarrow \infty} \left\| \int_0^t (\mathcal{A}_s^{(K)} \mathbf{T}_s + \mathbf{b}_s) ds - \int_0^t (\mathcal{A}_\infty^{(K)} \mathbf{T}_s + \mathbf{b}_\infty) ds \right\|_0 = 0. \quad (177)$$

Therefore,

$$\lim_{t \rightarrow \infty} \left\| \int_0^t (\mathcal{A}_\infty^{(K)} \mathbf{T}_s + \mathbf{b}_\infty) ds \right\|_0 = 0. \quad (178)$$

With $\tilde{\mathbf{u}}_\infty$ defined in (165), as $\mathbf{b}_\infty = \mathcal{A}_\infty^{(K)} \mathbf{u}_\infty$, there is

$$\lim_{t \rightarrow \infty} \left\| \mathcal{A}_\infty^{(K)} \left(\int_0^t \mathbf{T}_s ds - \mathbf{u}_\infty \right) \right\|_0 = 0. \quad (179)$$

Let $\xi^\parallel$ denote the component of a vector field $\xi \in \mathcal{V}(D)$ that is in the range of $\mathcal{A}_\infty^{(K)}$. In the ERM setting, $\mathcal{A}_\infty^{(K)}$ has a least nonzero eigenvalue that is positive, and hence the above implies that

$$\lim_{t \rightarrow \infty} \left\| \left(\int_0^t \mathbf{T}_s ds - \tilde{\mathbf{u}}_\infty \right)^\parallel \right\|_0 = 0 \quad (180)$$

or

$$\lim_{t \rightarrow \infty} \left\| \left(\int_0^t \mathbf{T}_s ds \right)^\parallel - \tilde{\mathbf{u}}_\infty \right\|_0 = 0 \quad (181)$$

and therefore, as $\text{Nul}(\mathcal{A}_\infty^{(K)}) = \text{Nul}(\mathcal{B}_\infty \mathcal{B}_\infty^\top) = \text{Nul}(\mathcal{B}_\infty^\top)$, it follows that

$$\lim_{t \rightarrow \infty} \left\| \mathcal{B}_\infty^\top \left(\int_0^t \mathbf{T}_s ds \right) - \mathcal{B}_\infty^\top \tilde{\mathbf{u}}_\infty \right\|_0 = 0. \quad (182)$$

Similar to (111), it can be shown that $\int_0^\infty \|\mathcal{B}_t - \mathcal{B}_\infty\|_0 dt < \infty$. Therefore, we have

$$\lim_{t \rightarrow \infty} \left\| \left(\int_0^t \mathcal{B}_s^\top \mathbf{T}_s ds \right) - \mathcal{B}_\infty^\top \tilde{\mathbf{u}}_\infty \right\|_0 = 0. \quad (183)$$

Now,

$$\begin{aligned} \int_0^t \langle \mathbf{T}_s, \mathcal{A}_s^{(K)} \mathbf{T}_s \rangle_0 ds &= \int_0^t \langle \mathcal{B}_s^\top \mathbf{T}_s, \mathcal{B}_s^\top \mathbf{T}_s \rangle_{\hat{\nu}, 0} ds \\ &\geq \left\langle \left(\int_0^t \mathcal{B}_s^\top \mathbf{T}_s ds \right), \left(\int_0^t \mathcal{B}_s^\top \mathbf{T}_s ds \right) \right\rangle_{\hat{\nu}, 0}. \end{aligned} \quad (184)$$

Hence,

$$\begin{aligned} \lim_{t \rightarrow \infty} \int_0^t \langle \mathbf{T}_s, \mathcal{A}_s^{(K)} \mathbf{T}_s \rangle_0 ds &\geq \lim_{t \rightarrow \infty} \left\langle \left(\int_0^t \mathcal{B}_s^\top \mathbf{T}_s ds \right), \left(\int_0^t \mathcal{B}_s^\top \mathbf{T}_s ds \right) \right\rangle_{\hat{\nu}, 0} \\ &= \langle \mathcal{B}_\infty^\top \tilde{\mathbf{u}}_\infty, \mathcal{B}_\infty^\top \tilde{\mathbf{u}}_\infty \rangle_{\hat{\nu}, 0} \\ &= \left\langle \mathcal{B}_\infty^\top \left(\mathcal{A}_\infty^{(K)} \right)^\dagger \mathbf{b}_\infty, \mathcal{B}_\infty^\top \left(\mathcal{A}_\infty^{(K)} \right)^\dagger \mathbf{b}_\infty \tilde{\mathbf{u}}_\infty \right\rangle_{\hat{\nu}, 0} \\ &= \left\langle \left(\mathcal{A}_\infty^{(K)} \right)^\dagger \mathbf{b}_\infty, \left(\mathcal{A}_\infty^{(K)} \right) \left(\mathcal{A}_\infty^{(K)} \right)^\dagger \mathbf{b}_\infty \right\rangle_0 \\ &= \left\langle \mathbf{b}_\infty, \left(\mathcal{A}_\infty^{(K)} \right)^\dagger \mathbf{b}_\infty \right\rangle_0. \end{aligned} \quad (185)$$

□

D.2.4 Proof of Lemma D.5

Since

$$\bar{g}_\infty(\mathbf{x}) = \int_D \varphi(\boldsymbol{\theta}, \mathbf{x}) \omega_0(d\boldsymbol{\theta}), \quad (186)$$

we know that when viewed as a L -dimensional random vector, $\bar{g}_\infty$ has the distribution

$$\bar{g}_\infty \sim \mathcal{N}(0, \bar{C}_\infty), \quad (187)$$

where

$$\begin{aligned} (\bar{C}_\infty)_{ij} &:= \mathbb{E}_0 [\bar{g}_\infty(\mathbf{x}_i) \bar{g}_\infty(\mathbf{x}_j)] \\ &= \int_D \varphi(\boldsymbol{\theta}, \mathbf{x}_i) \varphi(\boldsymbol{\theta}, \mathbf{x}_j) \mu_\infty(d\boldsymbol{\theta}) - \int_D \varphi(\boldsymbol{\theta}, \mathbf{x}_i) \mu_\infty(d\boldsymbol{\theta}) \int_D \varphi(\boldsymbol{\theta}', \mathbf{x}_j) \mu_\infty(d\boldsymbol{\theta}'), \end{aligned} \quad (188)$$

by the covariance of ω_0 , (32). Thus, we decompose $\bar{C}_\infty$ as $\bar{C}_\infty = \bar{C}_\infty^{(1)} - \bar{C}_\infty^{(2)}$, with

$$(\bar{C}_\infty^{(1)})_{ij} = \int_D \varphi(\boldsymbol{\theta}, \mathbf{x}_i) \varphi(\boldsymbol{\theta}, \mathbf{x}_j) \mu_\infty(d\boldsymbol{\theta}), \quad (189)$$

$$(\bar{C}_\infty^{(2)})_{ij} = \int_D \varphi(\boldsymbol{\theta}, \mathbf{x}_i) \mu_\infty(d\boldsymbol{\theta}) \int_D \varphi(\boldsymbol{\theta}', \mathbf{x}_j) \mu_\infty(d\boldsymbol{\theta}'). \quad (190)$$

Since $\bar{C}_\infty$ is PSD, its square root $(\bar{C}_\infty)^{1/2}$ is well-defined. By the property of multivariate Gaussian, we can write

$$\bar{g}_\infty \stackrel{d}{=} (\bar{C}_\infty)^{1/2} w, \quad (191)$$

where $\stackrel{d}{=}$ denotes equality in distribution, and $w \in \mathbb{R}^n$ follows the distribution

$$w \sim \mathcal{N}(0, \text{Id}_n). \quad (192)$$

This means that almost surely, $\bar{g}_\infty \in \text{Ran}((\bar{C}_\infty)^{1/2})$, and which would imply that $\bar{g}_\infty \in \text{Ran}(\bar{C}_\infty)$. This means that almost surely, we can write

$$\bar{g}_\infty = \bar{C}_\infty^{(1)} w^{(1)} - \bar{C}_\infty^{(2)} w^{(2)} \quad (193)$$

for some pair of $w^{(1)}, w^{(2)} \in \mathbb{R}^n$. Our goal is then to show that both $\bar{C}_\infty^{(1)} w^{(1)}$ and $\bar{C}_\infty^{(2)} w^{(2)}$ belong to $\text{Ran}(\mathcal{B}_\infty^\top)$. Here, under Assumption 2.1, since $\varphi(\boldsymbol{\theta}, \mathbf{x}) = c\hat{\varphi}(\mathbf{z}, \mathbf{x})$ when $\boldsymbol{\theta} = [c \quad \mathbf{z}]^\top$, there is

$$\nabla \varphi(\boldsymbol{\theta}, \mathbf{x}) = \begin{bmatrix} \hat{\varphi}(\mathbf{z}, \mathbf{x}) \\ c \nabla_{\mathbf{z}} \hat{\varphi}(\mathbf{z}, \mathbf{x}) \end{bmatrix}. \quad (194)$$

Therefore, first, we have

$$\begin{aligned} (\bar{C}_\infty^{(1)} w^{(1)})_i &= \int_D \varphi(\boldsymbol{\theta}, \mathbf{x}_i) \left(\sum_{j=1}^n \varphi(\boldsymbol{\theta}, \mathbf{x}_j) w_j^{(1)} \right) \mu_\infty(d\boldsymbol{\theta}) \\ &= \int_D \nabla \varphi(\boldsymbol{\theta}, \mathbf{x}_i)^\top \cdot \begin{bmatrix} c(\boldsymbol{\theta}) \left(\sum_{j=1}^n \varphi(\boldsymbol{\theta}, \mathbf{x}_j) w_j^{(1)} \right) \\ 0 \end{bmatrix} \mu_\infty(d\boldsymbol{\theta}) \\ &= \mathcal{B}_\infty^\top \boldsymbol{\xi}^{(1)}, \end{aligned} \quad (195)$$

with

$$\boldsymbol{\xi}(\boldsymbol{\theta})^{(1)} = \begin{bmatrix} c(\boldsymbol{\theta}) \left(\sum_{j=1}^n \varphi(\boldsymbol{\theta}, \mathbf{x}_j) w_j^{(1)} \right) \\ 0 \end{bmatrix}. \quad (196)$$

This means that $(\bar{C}_\infty^{(1)} w^{(1)}) \in \text{Ran}(\mathcal{B}_\infty^\top)$.

Second, there is

$$\begin{aligned} (\bar{C}_\infty^{(2)} w^{(2)})_i &= \left(\int_D \varphi(\boldsymbol{\theta}, \mathbf{x}_i) \mu_\infty(d\boldsymbol{\theta}) \right) \left(\sum_{j=1}^n w_j^{(2)} \int_D \varphi(\boldsymbol{\theta}', \mathbf{x}_j) \mu_\infty(d\boldsymbol{\theta}') \right) \\ &= \int_D \nabla \varphi(\boldsymbol{\theta}, \mathbf{x}_i)^\top \cdot \begin{bmatrix} c(\boldsymbol{\theta}) \left(\sum_{j=1}^n w_j^{(2)} \int_D \varphi(\boldsymbol{\theta}', \mathbf{x}_j) \mu_\infty(d\boldsymbol{\theta}') \right) \\ 0 \end{bmatrix} \mu_\infty(d\boldsymbol{\theta}) \\ &= \mathcal{B}_\infty^\top \boldsymbol{\xi}^{(2)}, \end{aligned} \quad (197)$$

with

$$\boldsymbol{\xi}(\boldsymbol{\theta})^{(2)} = \begin{bmatrix} c(\boldsymbol{\theta}) \left(\sum_{j=1}^n w_j^{(2)} \int_D \varphi(\boldsymbol{\theta}', \mathbf{x}_j) \mu_\infty(d\boldsymbol{\theta}') \right) \\ 0 \end{bmatrix} \quad (198)$$

This means that $(\bar{C}_\infty^{(2)} w^{(2)}) \in \text{Ran}(\mathcal{B}_\infty^\top)$. Hence the lemma is proved. $\square$

D.3 Proof of Theorem 3.6 (Under assumptions on the curvature in the long-time)

When the limiting measure μ_∞ does not necessarily interpolate the training data, such as in the regularized case, we have the following condition on T_t which guarantees that (46) holds:

Lemma D.7. *If*

$$\lim_{T \rightarrow \infty} \mathbb{E}_0 \int_0^T \int_D \langle T_t(\boldsymbol{\theta}), \nabla \nabla V(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mu_t) T_t(\boldsymbol{\theta}) \rangle \mu_0(d\boldsymbol{\theta}) dt \geq 0, \quad (199)$$

(including when this limit is $+\infty$) then (46) holds.

Proof of Lemma D.7: With $\mathfrak{D}_t$ defined in (101), for (46) to hold, it is sufficient to show that

$$\lim_{T \rightarrow \infty} \oint_0^T \mathfrak{D}_t dt \leq 0. \quad (200)$$

Recall from (107) that

$$\oint_0^T \mathfrak{D}_t dt = -\frac{1}{T} \|T_T\|_0^2 - 2 \oint_0^T \langle T_t, \mathcal{A}_t^{(V)} T_t \rangle_0 dt - \oint_0^T \langle T_t, \mathcal{A}_t^{(K)} T_t \rangle_0 dt. \quad (201)$$

Since $T_0 = 0$ and $\mathcal{A}_t^{(K)}$ is PSD, we see that the assumption (199) is sufficient. $\square$

Note that condition (199) is natural since we know from Proposition 2.7 that $\lim_{t \rightarrow \infty} \nabla \nabla V(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mu_t) = \nabla \nabla V(\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}), \mu_\infty)$ exists and is positive semidefinite μ_0 -almost surely. This lemma then allows us to prove Theorem 3.6: *Proof of Theorem 3.6:* Our goal is to verify (199) in order to apply Lemma D.7.

We first see that

$$\begin{aligned}
& \mathbb{E}_0 \int_D \langle \mathbf{T}_t(\boldsymbol{\theta}), \nabla \nabla V(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mu_t) \mathbf{T}_t(\boldsymbol{\theta}) \rangle \mu_0(d\boldsymbol{\theta}) \\
& \geq \mathbb{E}_0 \int_D \lambda_{\min}(\nabla \nabla V(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mu_t)) |\mathbf{T}_t(\boldsymbol{\theta})|^2 \mu_0(d\boldsymbol{\theta}) \\
& \geq \mathbb{E}_0 \int_D \min \{ \lambda_{\min}(\nabla \nabla V(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mu_t)), 0 \} |\mathbf{T}_t(\boldsymbol{\theta})|^2 \mu_0(d\boldsymbol{\theta}) \\
& = \int_D \min \{ \lambda_{\min}(\nabla \nabla V(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mu_t)), 0 \} (\mathbb{E}_0 |\mathbf{T}_t(\boldsymbol{\theta})|^2) \mu_0(d\boldsymbol{\theta}) \\
& \geq \int_D \min \{ \lambda_{\min}(\nabla \nabla V(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mu_t)), 0 \} \left(\sup_{\boldsymbol{\theta} \in \text{supp } \mu_0} \mathbb{E}_0 |\mathbf{T}_t(\boldsymbol{\theta})|^2 \right) \mu_0(d\boldsymbol{\theta}) \\
& \geq \|\mathbf{T}_t\|_{\text{sup}}^2 \left(\int_D \min \{ \lambda_{\min}(\nabla \nabla V(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mu_t)), 0 \} \mu_0(d\boldsymbol{\theta}) \right),
\end{aligned} \tag{202}$$

where we define, for $\boldsymbol{\xi} \in \mathcal{V}(D)$,

$$\|\boldsymbol{\xi}\|_{\text{sup}} := \sup_{\boldsymbol{\theta} \in \text{supp } \mu_0} \left(\mathbb{E}_0 |\boldsymbol{\xi}(\boldsymbol{\theta})|^2 \right)^{1/2}, \tag{203}$$

which is a norm on $\mathcal{V}(D)$.

Hence, if we assume that $\left| \int_D \min \{ \lambda_{\min}(\nabla \nabla V(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mu_t)), 0 \} \mu_0(d\boldsymbol{\theta}) \right|$ is small asymptotically, then what remains is to upper-bound $\|\mathbf{T}_t\|_{\text{sup}}$. Recall from (102) that the dynamics of $\mathbf{T}_t$ is governed by

$$\dot{\mathbf{T}}_t = -(\mathcal{A}_t^{(K)} + \mathcal{A}_t^{(V)})\mathbf{T}_t - \mathbf{b}_t, \tag{204}$$

Thus, in the $\|\cdot\|_{\text{sup}}$ norm defined above, we have

$$\begin{aligned}
\frac{d}{dt} \|\mathbf{T}_t\|_{\text{sup}} & \leq \| -(\mathcal{A}_t^{(K)} + \mathcal{A}_t^{(V)})\mathbf{T}_t - \mathbf{b}_t \|_{\text{sup}} \\
& \leq \|\mathcal{A}_t^{(K)}\mathbf{T}_t\|_{\text{sup}} + \|\mathcal{A}_t^{(V)}\mathbf{T}_t\|_{\text{sup}} + \|\mathbf{b}_t\|_{\text{sup}}.
\end{aligned} \tag{205}$$

We then want to bound the growth of $\|\mathbf{T}_t\|_{\text{sup}}$ by upper-bounding the RHS. Note that for $\boldsymbol{\xi} \in \mathcal{V}(D)$,

$$\begin{aligned}
\|\mathcal{A}_t^{(V)}\boldsymbol{\xi}\|_{\text{sup}}^2 & = \sup_{\boldsymbol{\theta} \in D} \mathbb{E}_0 |(\mathcal{A}_t^{(V)}\boldsymbol{\xi})(\boldsymbol{\theta})|^2 \\
& = \sup_{\boldsymbol{\theta} \in D} \mathbb{E}_0 |\nabla \nabla V(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mu_t) \boldsymbol{\xi}(\boldsymbol{\theta})|^2 \\
& \leq \sup_{\boldsymbol{\theta} \in D} |\nabla \nabla V(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mu_t)|^2 \mathbb{E}_0 |\boldsymbol{\xi}(\boldsymbol{\theta})|^2 \\
& \leq (C_{\nabla \nabla \varphi} C_{\varphi} + \lambda)^2 \sup_{\boldsymbol{\theta} \in D} \mathbb{E}_0 |\boldsymbol{\xi}(\boldsymbol{\theta})|^2 \\
& = (C_{\nabla \nabla \varphi} C_{\varphi} + \lambda)^2 \|\boldsymbol{\xi}\|_{\text{sup}}^2,
\end{aligned} \tag{206}$$

$$\begin{aligned}
\|\mathcal{A}_t^{(K)} \boldsymbol{\xi}\|_{\sup}^2 &= \sup_{\boldsymbol{\theta} \in D} \mathbb{E}_0 |(\mathcal{A}_t^{(K)} \boldsymbol{\xi})(\boldsymbol{\theta})|^2 \\
&= \sup_{\boldsymbol{\theta} \in D} \mathbb{E}_0 \left| \int_D \nabla' \nabla K(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \boldsymbol{\Theta}_t(\boldsymbol{\theta}')) \boldsymbol{\xi}(\boldsymbol{\theta}') \mu_0(d\boldsymbol{\theta}') \right|^2 \\
&\leq \sup_{\boldsymbol{\theta} \in D} \mathbb{E}_0 \int_D |\nabla' \nabla K(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \boldsymbol{\Theta}_t(\boldsymbol{\theta}'))|^2 |\boldsymbol{\xi}(\boldsymbol{\theta}')|^2 \mu_0(d\boldsymbol{\theta}') \\
&\leq \sup_{\boldsymbol{\theta} \in D} (C_{\nabla\varphi})^4 \int_D \mathbb{E}_0 |\boldsymbol{\xi}(\boldsymbol{\theta}')|^2 \mu_0(d\boldsymbol{\theta}') \\
&\leq (C_{\nabla\varphi})^4 \sup_{\boldsymbol{\theta}' \in D} \mathbb{E}_0 |\boldsymbol{\xi}(\boldsymbol{\theta}')|^2 \\
&= (C_{\nabla\varphi})^4 \|\boldsymbol{\xi}\|_{\sup}^2 .
\end{aligned} \tag{207}$$

Thus,

$$\|\mathcal{A}_t^{(K)} \mathbf{T}_t\|_{\sup} + \|\mathcal{A}_t^{(V)} \mathbf{T}_t\|_{\sup} \leq (C_{\nabla\varphi}^2 + C_{\nabla\varphi} C_{\varphi} + \lambda) \|\mathbf{T}_t\|_{\sup} . \tag{208}$$

To bound $\|\mathbf{b}_t\|_{\sup}$, we recall that

$$\begin{aligned}
\mathbf{b}_t(\boldsymbol{\theta}) &= \int_D \nabla K(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \boldsymbol{\Theta}_t(\boldsymbol{\theta}')) \omega_0(d\boldsymbol{\theta}') \\
&= \int_{\Omega} \nabla \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mathbf{x}) \bar{g}_t(\mathbf{x}) \hat{\nu}(d\mathbf{x}) ,
\end{aligned} \tag{209}$$

with

$$\bar{g}_t(\mathbf{x}) = \int_D \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mathbf{x}) \omega_0(d\boldsymbol{\theta}) . \tag{210}$$

This implies that $\forall \boldsymbol{\theta} \in \text{supp } \mu_0$,

$$|\mathbf{b}_t(\boldsymbol{\theta})| \leq \frac{1}{n} C_{\nabla\varphi} \sum_{l=1}^n |\bar{g}_t(\mathbf{x}_l)| \tag{211}$$

and so

$$\begin{aligned}
\mathbb{E}_0 |\mathbf{b}_t(\boldsymbol{\theta})|^2 &\leq C_{\nabla\varphi}^2 \mathbb{E}_0 \left(\frac{1}{n} \sum_{l=1}^n |\bar{g}_t(\mathbf{x}_l)| \right)^2 \\
&\leq C_{\nabla\varphi}^2 \mathbb{E}_0 \left(\frac{1}{n} \sum_{l=1}^n |\bar{g}_t(\mathbf{x}_l)|^2 \right) \\
&\leq C_{\nabla\varphi}^2 \frac{1}{n} \sum_{l=1}^n \mathbb{E}_0 |\bar{g}_t(\mathbf{x}_l)|^2 .
\end{aligned} \tag{212}$$

On the other hand, similar to (161), we have

$$\begin{aligned}
\mathbb{E}_0 |\bar{g}_t(\mathbf{x})|^2 &= \mathbb{E}_0 \left| \int_D \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mathbf{x}) \omega_0(d\boldsymbol{\theta}) \right|^2 \\
&= \int_D \left(\varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mathbf{x}) - \int_D \varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}'), \mathbf{x}) \mu_0(d\boldsymbol{\theta}') \right)^2 \mu_0(d\boldsymbol{\theta}) \\
&\leq \int_D |\varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mathbf{x})|^2 \mu_0(d\boldsymbol{\theta}) \\
&\leq (C_{\varphi})^2 ,
\end{aligned} \tag{213}$$

Thus, there is $\forall \boldsymbol{\theta} \in \text{supp } \mu_0$,

$$\mathbb{E}_0 |\mathbf{b}_t(\boldsymbol{\theta})|^2 \leq (C_{\nabla\varphi})^2 (C_\varphi)^2 \quad (214)$$

and so

$$\|\mathbf{b}_t\|_{\text{sup}} \leq C_{\nabla\varphi} C_\varphi . \quad (215)$$

Therefore, based on (205), we have

$$\frac{d}{dt} \|\mathbf{T}_t\|_{\text{sup}} \leq ((C_{\nabla\varphi})^2 + C_{\nabla\nabla\varphi} C_\varphi + \lambda) \|\mathbf{T}_t\|_{\text{sup}} + C_{\nabla\varphi} C_\varphi . \quad (216)$$

Since $\mathbf{T}_0 = 0$, we thus have

$$\begin{aligned} \|\mathbf{T}_t\|_{\text{sup}} &\leq C_{\nabla\varphi} C_\varphi \int_0^t e^{((C_{\nabla\varphi})^2 + C_{\nabla\nabla\varphi} C_\varphi + \lambda)(t-s)} ds \\ &= C_{\nabla\varphi} C_\varphi e^{((C_{\nabla\varphi})^2 + C_{\nabla\nabla\varphi} C_\varphi + \lambda)t} \int_0^t e^{-((C_{\nabla\varphi})^2 + C_{\nabla\nabla\varphi} C_\varphi + \lambda)s} ds \\ &\leq \frac{C_{\nabla\varphi} C_\varphi}{(C_{\nabla\varphi})^2 + C_{\nabla\nabla\varphi} C_\varphi + \lambda} e^{((C_{\nabla\varphi})^2 + C_{\nabla\nabla\varphi} C_\varphi + \lambda)t} \end{aligned} \quad (217)$$

Now, using (202), we see that in order for (199) to hold, it is sufficient to have

$$\lim_{t \rightarrow \infty} e^{((C_{\nabla\varphi})^2 + C_{\nabla\nabla\varphi} C_\varphi + \lambda)t} \left(\int_D \min \{ \lambda_{\min}(\nabla\nabla V(\boldsymbol{\theta}, \mu_t)), 0 \} \mu_0(d\boldsymbol{\theta}) \right) = 0 \quad (218)$$

and therefore sufficient to have

$$- \int_D \min \{ \lambda_{\min}(\nabla\nabla V(\boldsymbol{\theta}, \mu_t)), 0 \} \mu_0(d\boldsymbol{\theta}) \sim O \left(e^{-((C_{\nabla\varphi})^2 + C_{\nabla\nabla\varphi} C_\varphi + \lambda)t} \right) \quad (219)$$

□

To intuitively understand (50), note that we know from (23) in Proposition 2.7 that $\Lambda_t(\boldsymbol{\theta}) \rightarrow 0$ μ_0 -almost surely as $t \rightarrow \infty$. Condition (50) can therefore be satisfied by having $\Lambda_t(\boldsymbol{\theta})$ converge to zero sufficiently fast in the regions of D where it is negative, or having the measure of these regions with respect to μ_0 converge to zero sufficiently fast, or both.

D.4 Proof of Theorem 3.7 (Regularized case)

Recall from Proposition 3.3 that the dynamics of g_t is governed by

$$g_t(\mathbf{x}) + \int_0^t \int_\Omega \Gamma_{t,s}(\mathbf{x}, \mathbf{x}') g_s(\mathbf{x}') \hat{\nu}(d\mathbf{x}') ds = \bar{g}_t(\mathbf{x}) , \quad (220)$$

with

$$\Gamma_{t,s}(\mathbf{x}, \mathbf{x}') = \int_D \langle \nabla\varphi(\boldsymbol{\Theta}_t(\boldsymbol{\theta}), \mathbf{x}), J_{t,s}(\boldsymbol{\theta}) \nabla\varphi(\boldsymbol{\Theta}_s(\boldsymbol{\theta}), \mathbf{x}') \rangle \mu_0(d\boldsymbol{\theta}) , \quad (221)$$

with $J_{t,s}$ being the Jacobian of the flow $\boldsymbol{\Theta}_t$.

In the ERM setting, $\text{supp } \hat{\nu}$ is singular, thus we have $\hat{\nu}(d\mathbf{x}) = n^{-1} \sum_{l=1}^n \delta_{\mathbf{x}_l}(d\mathbf{x})$, where n is the total number of training data points. We define $\mathcal{W}_n(\Omega)$ together with the inner product $\langle \cdot, \cdot \rangle_{\hat{\nu},0}$ and the norm $\| \cdot \|_{\hat{\nu},0}$ as in Appendix A. We will also continue to consider g_t and $\bar{g}_t$ equivalently as n -dimensional vectors,

$$(g_t(\mathbf{x}_1) \ \cdots \ g_t(\mathbf{x}_n))^T , \quad (\bar{g}_t(\mathbf{x}_1) \ \cdots \ \bar{g}_t(\mathbf{x}_n))^T , \quad (222)$$

respectively. Thus, $\Gamma_{t,s}$ can also be represented by the $n \times n$ matrix

$$\begin{pmatrix} \Gamma_{t,s}(\mathbf{x}_1, \mathbf{x}_1) & \cdots & \Gamma_{t,s}(\mathbf{x}_1, \mathbf{x}_n) \\ \vdots & & \vdots \\ \Gamma_{t,s}(\mathbf{x}_n, \mathbf{x}_1) & \cdots & \Gamma_{t,s}(\mathbf{x}_n, \mathbf{x}_n) \end{pmatrix}. \quad (223)$$

Under such an abuse of notations, we can simplify (220) into

$$g_t + \int_0^t \Gamma_{t,s} g_s ds = \bar{g}_t. \quad (224)$$

Thus, the goal is to prove that

$$\lim_{t \rightarrow \infty} \sup \int_0^t \mathbb{E}_0 \|g_t\|_{\hat{\nu}}^2 dt \leq \mathbb{E}_0 \|\bar{g}_\infty\|_{\hat{\nu}}^2. \quad (225)$$

As in (43), we also define

$$\Gamma_{t-s}^\infty(\mathbf{x}, \mathbf{x}') = \int_D \langle \nabla \varphi(\boldsymbol{\theta}, \mathbf{x}), e^{-(t-s)\nabla \nabla V_\infty(\boldsymbol{\theta})} \nabla \varphi(\boldsymbol{\theta}, \mathbf{x}') \rangle \mu_\infty(d\boldsymbol{\theta}), \quad (226)$$

where for simplicity, we write $V_t(\cdot)$ for $V(\cdot, \mu_t)$ and $V_\infty(\cdot)$ for $V(\cdot, \mu_\infty)$. Then the heuristic argument outlined in Section 3.2 before Theorem 3.4 amounts to rewriting (224) as

$$g_t + \int_0^t \Gamma_{t-s}^\infty g_s ds = \bar{g}_t + \int_0^t (\Gamma_{t-s}^\infty - \Gamma_{t,s}) g_s ds \quad (227)$$

and then arguing that 1) Γ^∞ is a nonnegative convolution-type Volterra kernel, and 2) the second term on the RHS is small. Rigorously, we need to introduce an extra level of complication: for every $t_0 > 0$, we can rewrite (224) into

$$\begin{aligned} g_t &= \bar{g}_t - \int_{t_0}^t \Gamma_{t,s} g_s ds - \int_0^{t_0} \Gamma_{t,s} g_s ds \\ &= \bar{g}_t - \int_{t_0}^t \Gamma_{t-s}^\infty g_s ds + \int_{t_0}^t (\Gamma_{t-s}^\infty - \Gamma_{t,s}) g_s ds - \int_0^{t_0} \Gamma_{t,s} g_s ds. \end{aligned} \quad (228)$$

Then, for any $T > t_0$, by multiplying g_t and integrating from t_0 to T , we get

$$\begin{aligned} &\int_{t_0}^T \|g_t\|_{\hat{\nu}}^2 dt + \int_{t_0}^T \int_{t_0}^t \langle g_t, \Gamma_{t-s}^\infty g_s \rangle_{\hat{\nu}} ds dt \\ &\leq \int_{t_0}^T \langle g_t, \bar{g}_t \rangle_{\hat{\nu}} dt + \int_{t_0}^T \left\langle g_t, \int_{t_0}^t (\Gamma_{t-s}^\infty - \Gamma_{t,s}) g_s ds \right\rangle_{\hat{\nu}} dt + \int_{t_0}^T \left\langle g_t, \int_0^{t_0} \Gamma_{t,s} g_s ds \right\rangle_{\hat{\nu}} dt. \end{aligned} \quad (229)$$

Then firstly, the second term on the LHS is nonnegative because of the nonnegativity of Γ_t^∞ as a convolution-type Volterra kernel, as proven in Appendix D.1.

Hence, we have

$$\begin{aligned} \int_{t_0}^T \|g_t\|_{\hat{\nu}}^2 dt &\leq \int_{t_0}^T \langle g_t, \bar{g}_t \rangle_{\hat{\nu}} dt + \int_{t_0}^T \left\langle g_t, \int_{t_0}^t (\Gamma_{t-s}^\infty - \Gamma_{t,s}) g_s ds \right\rangle_{\hat{\nu}} dt \\ &\quad + \int_{t_0}^T \left\langle g_t, \int_0^{t_0} \Gamma_{t,s} g_s ds \right\rangle_{\hat{\nu}} dt. \end{aligned} \quad (230)$$

By Cauchy-Schwartz,

$$\int_{t_0}^T \langle g_t, \bar{g}_t \rangle_{\hat{\nu}} dt \leq \left(\int_{t_0}^T \|g_t\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}} \left(\int_{t_0}^T \|\bar{g}_t\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}}, \quad (231)$$

$$\begin{aligned} & \int_{t_0}^T \left\langle g_t, \int_{t_0}^t (\Gamma_{t-s}^\infty - \Gamma_{t,s}) g_s ds \right\rangle_{\hat{\nu}} dt \\ & \leq \left(\int_{t_0}^T \|g_t\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}} \left(\int_{t_0}^T \left\| \int_{t_0}^t (\Gamma_{t-s}^\infty - \Gamma_{t,s}) g_s ds \right\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}} \\ & \leq \left(\int_{t_0}^T \|g_t\|_{\hat{\nu}}^2 dt \right) \left(\int_{t_0}^T \int_{t_0}^t \|\Gamma_{t-s}^\infty - \Gamma_{t,s}\|_{\hat{\nu}}^2 ds dt \right)^{\frac{1}{2}}, \end{aligned} \quad (232)$$

and

$$\begin{aligned} & \int_{t_0}^T \left\langle g_t, \int_0^{t_0} \Gamma_{t,s} g_s ds \right\rangle_{\hat{\nu}} dt \\ & \leq \left(\int_{t_0}^T \|g_t\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}} \left(\int_{t_0}^T \left\| \int_0^{t_0} \Gamma_{t,s} g_s ds \right\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}} \\ & \leq \left(\int_{t_0}^T \|g_t\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}} \left(\int_{t_0}^T \left(\int_0^{t_0} \|\Gamma_{t,s}\|_{\hat{\nu}}^2 ds \right) \left(\int_0^{t_0} \|g_s\|_{\hat{\nu}}^2 ds \right) dt \right)^{\frac{1}{2}} \\ & \leq \left(\int_{t_0}^T \|g_t\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}} \left(\int_0^{t_0} \|g_t\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}} \left(\int_{t_0}^T \int_0^{t_0} \|\Gamma_{t,s}\|_{\hat{\nu}}^2 ds dt \right)^{\frac{1}{2}}. \end{aligned} \quad (233)$$

Therefore, putting everything together, we have

$$\begin{aligned} \left(\int_{t_0}^T \|g_t\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}} & \leq \left(\int_{t_0}^T \|\bar{g}_t\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}} \\ & + \left(\int_{t_0}^T \|g_t\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}} \left(\int_{t_0}^T \int_{t_0}^t \|\Gamma_{t-s}^\infty - \Gamma_{t,s}\|_{\hat{\nu}}^2 ds dt \right)^{\frac{1}{2}} \\ & + \left(\int_0^{t_0} \|g_t\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}} \left(\int_{t_0}^T \int_0^{t_0} \|\Gamma_{t,s}\|_{\hat{\nu}}^2 ds dt \right)^{\frac{1}{2}}, \end{aligned} \quad (234)$$

and hence, using $f_a^b \cdot dt$ to denote the averaged integral $\frac{1}{b-a} \int_a^b \cdot dt$,

$$\begin{aligned} \left(f_{t_0}^T \|g_t\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}} & \leq \left(f_{t_0}^T \|\bar{g}_t\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}} \\ & + \left(f_{t_0}^T \|g_t\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}} \left(f_{t_0}^T \int_{t_0}^t \|\Gamma_{t-s}^\infty - \Gamma_{t,s}\|_{\hat{\nu}}^2 ds dt \right)^{\frac{1}{2}} \\ & + \left(\int_0^{t_0} \|g_t\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}} \left(f_{t_0}^T \int_0^{t_0} \|\Gamma_{t,s}\|_{\hat{\nu}}^2 ds dt \right)^{\frac{1}{2}}, \end{aligned} \quad (235)$$

or

$$\begin{aligned} & \left(1 - \left[\int_{t_0}^T \int_{t_0}^t \|\Gamma_{t-s}^\infty - \Gamma_{t,s}\|_{\hat{\nu}}^2 ds dt \right]^{\frac{1}{2}} \right) \left(\int_{t_0}^T \|g_t\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}} \\ & \leq \left(\int_{t_0}^T \|\bar{g}_t\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}} + \left(\int_0^{t_0} \|g_t\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}} \left(\int_{t_0}^T \int_0^{t_0} \|\Gamma_{t,s}\|_{\hat{\nu}}^2 ds dt \right)^{\frac{1}{2}} . \end{aligned} \quad (236)$$

Lemma D.8. *Under all assumptions in Theorem 3.7 except for (52) being replaced by a weaker condition,*

$$\int_0^\infty \int_D (|\Theta_t(\theta) - \Theta_\infty(\theta)| + |U_t(\theta)|^2) e^{C_1(U_t(\theta) + \bar{U}_t)} \mu_0(d\theta) dt < \infty , \quad (237)$$

we have

$$\lim_{t_0 \rightarrow \infty} \int_{t_0}^\infty \int_{t_0}^t \|\Gamma_{t-s}^\infty - \Gamma_{t,s}\|_{\hat{\nu}}^2 ds dt = 0 \quad (238)$$

and $\forall t_0 > 0$,

$$\lim_{T \rightarrow \infty} \int_{t_0}^T \int_0^{t_0} \|\Gamma_{t,s}\|^2 ds dt = 0 . \quad (239)$$

We will prove in Appendix D.4.2 that (52) indeed implies (237).

The lemma will be proved in Appendix D.4.1, and let us first proceed with the proof of the theorem assuming this lemma. Suppose for contradiction that (225) does not hold, meaning that

$$\lim_{T \rightarrow \infty} \sup \left(\int_0^T \|g_t\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}} = \|\bar{g}_\infty\|_{\hat{\nu}} + \epsilon \quad (240)$$

for some $\epsilon > 0$. We will select a pair of t_0 and T for which the inequality (236) cannot be satisfied. Firstly, by the convergence of $\bar{g}_t$ to $\bar{g}_\infty$, $\exists t_a > 0$ such that $\forall t_1, t_2 > t_a$,

$$\left(\int_{t_1}^{t_2} \|\bar{g}_t\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}} \leq \|\bar{g}_\infty\|_{\hat{\nu}} + \frac{1}{6}\epsilon . \quad (241)$$

Secondly, by our assumption (240) and the first part of Lemma D.8, $\exists t_0 > t_a$ such that both

$$\left(\int_0^{t_0} \|g_t\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}} \leq \|\bar{g}_\infty\|_{\hat{\nu}} + 2\epsilon \quad (242)$$

and

$$\int_{t_0}^\infty \int_{t_0}^t \|\Gamma_{t-s}^\infty - \Gamma_{t,s}\|_{\hat{\nu}}^2 ds dt < \frac{\epsilon}{6\|\bar{g}_\infty\|_{\hat{\nu}} + 3\epsilon} \quad (243)$$

are satisfied. In particular, (242) implies

$$\left(\int_0^{t_0} |g_t|^2 dt \right)^{\frac{1}{2}} \leq t_0^{\frac{1}{2}} \cdot (\|\bar{g}_\infty\|_{\hat{\nu}} + 2\epsilon) \quad (244)$$

Let

$$\delta = \left(\frac{\epsilon}{6t_0^{\frac{1}{2}} \cdot (\|\bar{g}_\infty\|_{\hat{\nu}} + 2\epsilon)} \right)^2 > 0 . \quad (245)$$

By the second part of Lemma D.8, $\exists t_b > t_0$ such that $\forall T > t_b$,

$$\int_{t_0}^T \int_0^{t_0} \|\Gamma_{t,s}\|^2 ds dt < \delta \quad (246)$$

so that the last term in (236) satisfies

$$\left(\int_0^{t_0} \|g_t\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}} \left(\int_{t_0}^T \int_0^{t_0} \|\Gamma_{t,s}\|^2 ds dt \right)^{\frac{1}{2}} < \frac{1}{6}\epsilon \quad (247)$$

By our assumption (240), we can choose a $T > t_b$ such that

$$\left(\int_0^T \|g_t\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}} \geq \|\bar{g}_\infty\|_{\hat{\nu}} + \frac{2}{3}\epsilon. \quad (248)$$

Since

$$\left(\int_0^{t_0} \|g_t\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}} \leq \|\bar{g}_\infty\|_{\hat{\nu}} + 2\epsilon, \quad (249)$$

we can assume without loss of generality that $\frac{T}{t_0}$ is large enough so that

$$\left(\int_{t_0}^T \|g_t\|_{\hat{\nu}}^2 dt \right)^{\frac{1}{2}} \geq \|\bar{g}_\infty\|_{\hat{\nu}} + \frac{1}{2}\epsilon. \quad (250)$$

Thus, back to the inequality (236), the LHS is strictly lower-bounded by

$$(\|\bar{g}_\infty\|_{\hat{\nu}} + \frac{1}{2}\epsilon) \left(1 - \frac{\epsilon}{6\|\bar{g}_\infty\|_{\hat{\nu}} + 3\epsilon} \right) = \|\bar{g}_\infty\|_{\hat{\nu}} + \frac{1}{3}\epsilon, \quad (251)$$

whereas the RHS is strictly upper-bounded by

$$\|\bar{g}_\infty\|_{\hat{\nu}} + \frac{1}{6}\epsilon + \frac{1}{6}\epsilon = \|\bar{g}_\infty\|_{\hat{\nu}} + \frac{1}{3}\epsilon. \quad (252)$$

This gives contradiction and we are done with the proof of Theorem 3.7. $\square$

D.4.1 Proof of Lemma D.8

It remains to prove Lemma D.8. To do so we will need an auxiliary result, that we state and prove first:

Lemma D.9. *Let $\Delta\Gamma_{t,s} := \Gamma_{t,s} - \Gamma_{t-s}^\infty$. If $\nabla\nabla V$ is uniformly positive definite with eigenvalues lower-bounded by λ , then there exists constants C and C' whose values depend on $|D'|$, C_φ , $C_{\nabla\varphi}$, $C_{\nabla\nabla\varphi}$, and $L_{\nabla\nabla\varphi}$ such that*

$$\|\Delta\Gamma_{t,s}\|_{\hat{\nu}} \leq C e^{-\lambda(t-s)} \int_D \left(|\Delta\Theta_t(\theta)| + (|\Delta\Theta_s(\theta)| + U_s(\theta)) e^{C'(U_s(\theta) + \bar{U}_s)} \right) \mu_0(d\theta) \quad (253)$$

where $\Delta\Theta_t(\theta) = \Theta_t(\theta) - \Theta_\infty(\theta)$.

Proof of Lemma D.9: To bound $\|\Delta\Gamma_{t,s}\|_{\hat{\nu}}$, we bound $\|\Delta\Gamma_{t,s}\eta\|_{\hat{\nu}}$ for $\eta \in \mathbb{R}^n$. Note that $\Delta\Gamma_{t,s}\eta$ can be obtained in the following way. Consider the two systems

$$\begin{cases} \frac{d}{dt} \xi_t(\theta) = -\nabla\nabla V_t(\Theta_t(\theta)) \xi_t(\theta) \\ \xi_s(\theta) = \int_\Omega \nabla\varphi(\Theta_s(\theta), x') \eta(x') \hat{\nu}(dx') \end{cases} \quad (254)$$

$$\begin{cases} \frac{d}{dt}\xi'_t(\theta) = -\nabla\nabla V_\infty(\Theta_\infty(\theta))\xi'_t(\theta) \\ \xi'_s(\theta) = \int_\Omega \nabla\varphi(\Theta_\infty(\theta), x')\eta(x')\hat{\nu}(dx') \end{cases} \quad (255)$$

Then there is

$$\begin{aligned} (\Gamma_{t,s}\eta)(x) &= \int_D \nabla\varphi(\Theta_t(\theta), x) \cdot \xi_t(\theta) \mu_0(d\theta) \\ (\Gamma_{t-s}^\infty\eta)(x) &= \int_D \nabla\varphi(\Theta_\infty(\theta), x) \cdot \xi'_t(\theta) \mu_0(d\theta) \end{aligned} \quad (256)$$

and hence

$$\begin{aligned} (\Delta\Gamma_{t,s}\eta)(x) &= \int_D \nabla\varphi(\Theta_t(\theta), x) \xi_t(\theta) \mu_0(d\theta) - \int_D \nabla\varphi(\Theta_\infty(\theta), x) \xi'_t(\theta) \mu_0(d\theta) \\ &= \int_D \nabla\varphi(\Theta_t(\theta), x) \cdot (\xi_t(\theta) - \xi'_t(\theta)) \mu_0(d\theta) \\ &\quad + \int_D (\nabla\varphi(\Theta_t(\theta), x) - \nabla\varphi(\Theta_\infty(\theta), x)) \cdot \xi'_t(\theta) \mu_0(d\theta). \end{aligned} \quad (257)$$

We will first try to bound $\xi_t(\theta) - \xi'_t(\theta)$ as a function of η . Define $\Delta\xi_t(\theta) = \xi_t(\theta) - \xi'_t(\theta)$. Then

$$\begin{aligned} \frac{d}{dr}\Delta\xi_r(\theta) &= -(\nabla\nabla V_r(\Theta_r(\theta)) - \nabla\nabla V_\infty(\Theta_\infty(\theta)))\xi_r - \nabla\nabla V_\infty(\Theta_\infty(\theta))\Delta\xi_r(\theta) \\ &= -\nabla\nabla V_\infty(\Theta_\infty(\theta))\Delta\xi_r(\theta) \\ &\quad - (\nabla\nabla V_r(\Theta_r(\theta)) - \nabla\nabla V_\infty(\Theta_\infty(\theta)))\xi'_r \\ &\quad - (\nabla\nabla V_r(\Theta_r(\theta)) - \nabla\nabla V_\infty(\Theta_\infty(\theta)))\Delta\xi_r. \end{aligned} \quad (258)$$

Thus,

$$\begin{aligned} \Delta\xi_t(\theta) &= e^{-(t-s)\nabla\nabla V_\infty(\Theta_\infty(\theta))}\Delta\xi_s(\theta) \\ &\quad + \int_s^t e^{-(t-r)\nabla\nabla V_\infty(\Theta_\infty(\theta))}(\nabla\nabla V_r(\Theta_r(\theta)) - \nabla\nabla V_\infty(\Theta_\infty(\theta)))\xi'_r(\theta)dr \\ &\quad + \int_s^t e^{-(t-r)\nabla\nabla V_\infty(\Theta_\infty(\theta))}(\nabla\nabla V_r(\Theta_r(\theta)) - \nabla\nabla V_\infty(\Theta_\infty(\theta)))\Delta\xi_r(\theta)dr. \end{aligned} \quad (259)$$

Since $\nabla\nabla V_\infty(\Theta_\infty(\theta)) - \lambda I_d$ is positive semidefinite, we first have

$$|\xi'_r(\theta)| \leq e^{-\lambda(r-s)}|\xi'_s(\theta)| \quad (260)$$

as well as

$$\begin{aligned} |\Delta\xi_t(\theta)| &\leq e^{-\lambda(t-s)}|\Delta\xi_s(\theta)| \\ &\quad + \int_s^t e^{-\lambda(t-r)}\|\nabla\nabla V_r(\Theta_r(\theta)) - \nabla\nabla V_\infty(\Theta_\infty(\theta))\||\xi'_r(\theta)|dr \\ &\quad + \int_s^t e^{-\lambda(t-r)}\|\nabla\nabla V_r(\Theta_r(\theta)) - \nabla\nabla V_\infty(\Theta_\infty(\theta))\||\Delta\xi_r(\theta)|dr \\ &\leq e^{-\lambda(t-s)}|\Delta\xi_s(\theta)| \\ &\quad + \int_s^t e^{-\lambda(t-s)}\|\nabla\nabla V_r(\Theta_r(\theta)) - \nabla\nabla V_\infty(\Theta_\infty(\theta))\||\xi'_s(\theta)|dr \\ &\quad + \int_s^t e^{-\lambda(t-r)}\|\nabla\nabla V_r(\Theta_r(\theta)) - \nabla\nabla V_\infty(\Theta_\infty(\theta))\||\Delta\xi_r(\theta)|dr. \end{aligned} \quad (261)$$

To prepare for an application of Gronwall's inequality, we introduce a change-of-variable by defining, for $r \in [s, t]$,

$$\overline{\Delta \xi}_r(\theta) = e^{\lambda(t-s)} \Delta \xi_r(\theta). \quad (262)$$

Then we can rewrite the equation above as

$$\begin{aligned} |\overline{\Delta \xi}_t(\theta)| &= e^{\lambda(r-s)} |\Delta \xi_t(\theta)| \\ &\leq |\Delta \xi_s(\theta)| + \int_s^t \|\nabla \nabla V_r(\Theta_r(\theta)) - \nabla \nabla V_\infty(\Theta_\infty(\theta))\| \|\xi'_s(\theta)\| dr \\ &\quad + \int_s^t e^{\lambda(r-s)} \|\nabla \nabla V_r(\Theta_r(\theta)) - \nabla \nabla V_\infty(\Theta_\infty(\theta))\| |\Delta \xi_r(\theta)| dr \\ &\leq |\overline{\Delta \xi}_s(\theta)| + \int_s^t \|\nabla \nabla V_r(\Theta_r(\theta)) - \nabla \nabla V_\infty(\Theta_\infty(\theta))\| \|\xi'_s(\theta)\| dr \\ &\quad + \int_s^t \|\nabla \nabla V_r(\Theta_r(\theta)) - \nabla \nabla V_\infty(\Theta_\infty(\theta))\| |\overline{\Delta \xi}_r(\theta)| dr. \end{aligned} \quad (263)$$

Thus, by Gronwall's inequality,

$$\begin{aligned} |\overline{\Delta \xi}_t(\theta)| &\leq \left(|\overline{\Delta \xi}_s(\theta)| + \int_s^t \|\nabla \nabla V_r(\Theta_r(\theta)) - \nabla \nabla V_\infty(\Theta_\infty(\theta))\| \|\xi'_s(\theta)\| dr \right) \\ &\quad \times e^{\int_s^t \|\nabla \nabla V_r(\Theta_r(\theta)) - \nabla \nabla V_\infty(\Theta_\infty(\theta))\| dr}, \end{aligned} \quad (264)$$

or, back in the original variable that we are interested in,

$$\begin{aligned} |\Delta \xi_t(\theta)| &\leq \left(|\Delta \xi_s(\theta)| + \int_s^t \|\nabla \nabla V_r(\Theta_r(\theta)) - \nabla \nabla V_\infty(\Theta_\infty(\theta))\| \|\xi'_s(\theta)\| dr \right) \\ &\quad \times e^{-\lambda(t-s) + \int_s^t \|\nabla \nabla V_r(\Theta_r(\theta)) - \nabla \nabla V_\infty(\Theta_\infty(\theta))\| dr}. \end{aligned} \quad (265)$$

Now, we have

$$\begin{aligned} &|\Delta \Gamma_{t,s} \eta(\mathbf{x})| \\ &\leq \left\| \int_D \nabla \varphi(\Theta_t(\theta), \mathbf{x})^\top \cdot \Delta \xi_t(\theta) \mu_0(d\theta) \right\|_{\hat{\nu}} \\ &\quad + \left\| \int_D (\nabla \varphi(\Theta_t(\theta), \mathbf{x}) - \nabla \varphi(\Theta_\infty(\theta), \mathbf{x}))^\top \xi'_t(\theta) \mu_0(d\theta) \right\|_{\hat{\nu}} \\ &\leq C_{\nabla \varphi} \int_D |\Delta \xi_t(\theta)| \mu_0(d\theta) + C_{\nabla \nabla \varphi} \int_D |\Delta \Theta_t(\theta)| \|\xi'_t(\theta)\| \mu_0(d\theta) \\ &\leq C_{\nabla \varphi} e^{-\lambda(t-s)} \int_D \left(|\Delta \xi_s(\theta)| + \int_s^t \|\nabla \nabla V_r(\Theta_r(\theta)) - \nabla \nabla V_\infty(\Theta_\infty(\theta))\| \|\xi'_s(\theta)\| dr \right) \\ &\quad e^{\int_s^t \|\nabla \nabla V_r(\Theta_r(\theta)) - \nabla \nabla V_\infty(\Theta_\infty(\theta))\| dr} \mu_0(d\theta) \\ &\quad + C_{\nabla \nabla \varphi} e^{-\lambda(t-s)} \int_D |\Delta \Theta_t(\theta)| \|\xi'_s(\theta)\| \mu_0(d\theta). \end{aligned} \quad (266)$$

Note that we have,

$$|\xi'_s(\theta)| = \left| \int_\Omega \nabla \varphi(\Theta_\infty(\theta), \mathbf{x}') \eta(\mathbf{x}') \hat{\nu}(d\mathbf{x}') \right| \leq C_{\nabla \varphi} \sup_{1 \leq p \leq P} |\eta(\mathbf{x}_p)| \leq P^{\frac{1}{2}} C_{\nabla \varphi} \|\eta\|_{\hat{\nu}}, \quad (267)$$

$$\begin{aligned}
|\Delta \xi_s(\theta)| &= \left| \int_{\Omega} (\nabla \varphi(\Theta_s(\theta), x) - \nabla \varphi(\Theta_{\infty}(\theta), x)) \eta(x') \hat{\nu}(dx') \right| \\
&\leq \int_{\Omega} |\nabla \varphi(\Theta_s(\theta), x') - \nabla \varphi(\Theta_{\infty}(\theta), x')| |\eta(x')| \hat{\nu}(dx') \\
&\leq n^{\frac{1}{2}} C_{\nabla \nabla \varphi} |\Delta \Theta_s(\theta)| \|\eta\|_{\hat{\nu}}
\end{aligned} \tag{268}$$

and, since $\nabla \nabla V_r(\theta) = \int_{\Omega} \nabla \nabla \varphi(\theta, x)(f_r(x) - f_*(x)) \hat{\nu}(dx)$ and $\nabla \nabla V_{\infty}(\theta) = \int_{\Omega} \nabla \nabla \varphi(\theta, x)(f_{\infty}(x) - f_*(x)) \hat{\nu}(dx)$,

$$\begin{aligned}
&\|\nabla \nabla V_r(\Theta_r(\theta)) - \nabla \nabla V_{\infty}(\Theta_{\infty}(\theta))\| \\
&\leq \|\nabla \nabla V_r(\Theta_r(\theta)) - \nabla \nabla V_r(\Theta_{\infty}(\theta))\| \\
&\quad + \|\nabla \nabla V_r(\Theta_{\infty}(\theta)) - \nabla \nabla V_{\infty}(\Theta_{\infty}(\theta))\| \\
&\leq L_{\nabla \nabla \varphi} C_{\varphi} |\Delta \Theta_r(\theta)| + C_{\nabla \nabla \varphi} \|\Delta f_r\|_{\hat{\nu}, \infty},
\end{aligned} \tag{269}$$

where we use $\|f\|_{\hat{\nu}, \infty}$ to denote $\sup_{x \in \text{supp } \hat{\nu}} |f(x)|$ and we defined $\Delta f_t = f_t - f_{\infty}$.

As a result, we have

$$\begin{aligned}
&\|\Delta \Gamma_{t,s}\|_{\hat{\nu}} \\
&\leq \|\eta\|_{\hat{\nu}}^{-1} \|\Delta \Gamma_{t,s} \eta\|_{\hat{\nu}} \\
&\leq C_{\nabla \varphi} e^{-\lambda(t-s)} \int_D \left(C_{\nabla \nabla \varphi} |\Delta \Theta_s(\theta)| + \int_s^t (L_{\nabla \nabla \varphi} C_{\varphi} |\Delta \Theta_r(\theta)| + C_{\nabla \nabla \varphi} \|\Delta f_r\|_{\hat{\nu}, \infty}) C_{\nabla \varphi} dr \right) \\
&\quad \times e^{\int_s^t L_{\nabla \nabla \varphi} C_{\varphi} |\Delta \Theta_r(\theta)| + C_{\nabla \nabla \varphi} \|\Delta f_r\|_{\hat{\nu}, \infty} dr} \mu_0(d\theta) \\
&\quad + C_{\nabla \nabla \varphi} e^{-\lambda(t-s)} \int_D C_{\nabla \varphi} |\Delta \Theta_t(\theta)| \mu_0(d\theta).
\end{aligned} \tag{270}$$

Therefore, using C_0, C_1 , etc. to represent constants that depend on $C_{\varphi}, C_{\nabla \varphi}, C_{\nabla \nabla \varphi}, C_{\nabla \nabla \varphi}$ and $L_{\nabla \nabla \varphi}$, we have

$$\begin{aligned}
&\|\Delta \Gamma_{t,s}\|_{\hat{\nu}} \\
&\leq C_0 e^{-\lambda(t-s)} \left(\int_D |\Delta \Theta_t(\theta)| \mu_0(d\theta) + \int_D |\Delta \Theta_s(\theta)| e^{C_1 \int_s^t |\Delta \Theta_r(\theta)| + \|\Delta f_r\|_{\hat{\nu}, \infty} dr} \mu_0(d\theta) \right. \\
&\quad \left. + \int_D \left(\int_s^t |\Delta \Theta_r(\theta)| + \|f_r - f_{\infty}\|_{\hat{\nu}, \infty} dr \right) e^{C_1 \int_s^t |\Delta \Theta_r(\theta)| + \|\Delta f_r\|_{\hat{\nu}, \infty} dr} \mu_0(d\theta) \right).
\end{aligned} \tag{271}$$

Note that $\|\Delta f_r\|_{\hat{\nu}, \infty}$ can be further upper-bounded by $C_{\varphi} \int_D |\Delta \Theta_r(\theta)| \mu_0(d\theta)$. Furthermore, defining

$$\overline{\Delta \Theta_t} = \int_D |\Delta \Theta_t(\theta)| \mu_0(d\theta) \tag{272}$$

we can write the bound above as

$$\begin{aligned}
\|\Delta \Gamma_{t,s}\|_{\hat{\nu}} &\leq C_0 e^{-\lambda(t-s)} \left(\int_D |\Delta \Theta_t(\theta)| \mu_0(d\theta) + \int_D |\Delta \Theta_s(\theta)| e^{C_1 \int_s^t |\Delta \Theta_r(\theta)| + \overline{\Delta \Theta_r} dr} \mu_0(d\theta) \right. \\
&\quad \left. + \int_D \left(\int_s^t |\Delta \Theta_r(\theta)| + \overline{\Delta \Theta_r} dr \right) e^{C_1 \int_s^t |\Delta \Theta_r(\theta)| + \overline{\Delta \Theta_r} dr} \mu_0(d\theta) \right).
\end{aligned} \tag{273}$$

Finally, let

$$U_t(\theta) = \int_t^{\infty} |\Delta \Theta_t(\theta)| dt \tag{274}$$

and

$$\bar{U}_t = \int_D U_t(\boldsymbol{\theta}) \mu_0(d\boldsymbol{\theta}) = \int_t^\infty \overline{\Delta \boldsymbol{\Theta}_t} dt. \quad (275)$$

Then there is

$$\begin{aligned} \|\Delta \Gamma_{t,s}\|_{\hat{\nu}} &\leq C_0 e^{-\lambda(t-s)} \int_D \left(|\Delta \boldsymbol{\Theta}_t(\boldsymbol{\theta})| + (|\Delta \boldsymbol{\Theta}_s(\boldsymbol{\theta})| + U_s(\boldsymbol{\theta}) + \bar{U}_s) e^{C_1(U_s(\boldsymbol{\theta}) + \bar{U}_s)} \right) \mu_0(d\boldsymbol{\theta}) \\ &\leq 2C_0 e^{-\lambda(t-s)} \int_D \left(|\Delta \boldsymbol{\Theta}_t(\boldsymbol{\theta})| + (|\Delta \boldsymbol{\Theta}_s(\boldsymbol{\theta})| + U_s(\boldsymbol{\theta})) e^{C_1(U_s(\boldsymbol{\theta}) + \bar{U}_s)} \right) \mu_0(d\boldsymbol{\theta}). \end{aligned} \quad (276)$$

(End of the proof of Lemma D.9.) $\square$

Proof of Lemma D.8: Lemma D.9 entails that, $\exists C, C' > 0$ such that

$$\begin{aligned} \|\Delta \Gamma_{t,s}\|_{\hat{\nu}}^2 &\leq C e^{-2\lambda(t-s)} \left(\int_D \left(|\Delta \boldsymbol{\Theta}_t(\boldsymbol{\theta})| + (|\Delta \boldsymbol{\Theta}_s(\boldsymbol{\theta})| + U_s(\boldsymbol{\theta})) e^{C'(U_s(\boldsymbol{\theta}) + \bar{U}_s)} \right) \mu_0(d\boldsymbol{\theta}) \right)^2 \\ &\leq 4C e^{-2\lambda(t-s)} \int_D |\Delta \boldsymbol{\Theta}_t(\boldsymbol{\theta})|^2 + (|\Delta \boldsymbol{\Theta}_s(\boldsymbol{\theta})|^2 + U_s(\boldsymbol{\theta})^2) e^{2C'(U_s(\boldsymbol{\theta}) + \bar{U}_s)} \mu_0(d\boldsymbol{\theta}) \\ &\leq 4C |D'| e^{-2\lambda(t-s)} \int_D |\Delta \boldsymbol{\Theta}_t(\boldsymbol{\theta})| + (|\Delta \boldsymbol{\Theta}_s(\boldsymbol{\theta})| + U_s(\boldsymbol{\theta})^2) e^{2C'(U_s(\boldsymbol{\theta}) + \bar{U}_s)} \mu_0(d\boldsymbol{\theta}), \end{aligned} \quad (277)$$

where for the last inequality, we assume that $|D'| \geq 1$ (or, to accommodate the more general case, just replace $|D'|$ by $\max\{|D'|, 1\}$).

To prove Lemma D.8, the first goal is to show

$$\lim_{t_0 \rightarrow \infty} \int_{t_0}^\infty \int_{t_0}^t \|\Delta \Gamma_{t,s}\|_{\hat{\nu}}^2 ds dt = 0. \quad (278)$$

There is

$$\begin{aligned} &\int_{t_0}^\infty \int_{t_0}^t \|\Delta \Gamma_{t,s}\|_{\hat{\nu}}^2 ds dt \\ &\leq 4C |D'| \int_D \int_{t_0}^\infty \int_{t_0}^t e^{-2\lambda(t-s)} \left(|\Delta \boldsymbol{\Theta}_t(\boldsymbol{\theta})| + (|\Delta \boldsymbol{\Theta}_s(\boldsymbol{\theta})| + U_s(\boldsymbol{\theta})^2) e^{2C'(U_s(\boldsymbol{\theta}) + \bar{U}_s)} \right) ds dt \mu_0(d\boldsymbol{\theta}) \\ &\leq 4C |D'| \int_D \left(\int_{t_0}^\infty \left(\int_{t_0}^t e^{-2\lambda(t-s)} ds \right) |\Delta \boldsymbol{\Theta}_t(\boldsymbol{\theta})| dt \right. \\ &\quad \left. + \int_{t_0}^\infty \left(\int_s^\infty e^{-2\lambda(t-s)} dt \right) (|\Delta \boldsymbol{\Theta}_s(\boldsymbol{\theta})| + U_s(\boldsymbol{\theta})^2) e^{2C'(U_s(\boldsymbol{\theta}) + \bar{U}_s)} ds \right) \mu_0(d\boldsymbol{\theta}) \\ &\leq 2C |D'| \lambda^{-1} \int_D \left(\int_{t_0}^\infty |\Delta \boldsymbol{\Theta}_t(\boldsymbol{\theta})| dt + \int_{t_0}^\infty (|\Delta \boldsymbol{\Theta}_s(\boldsymbol{\theta})| + U_s(\boldsymbol{\theta})^2) e^{2C'(U_s(\boldsymbol{\theta}) + \bar{U}_s)} ds \right) \mu_0(d\boldsymbol{\theta}) \\ &\leq 4C |D'| \lambda^{-1} \int_D \int_{t_0}^\infty (|\Delta \boldsymbol{\Theta}_s(\boldsymbol{\theta})| + U_s(\boldsymbol{\theta})^2) e^{2C'(U_s(\boldsymbol{\theta}) + \bar{U}_s)} ds \mu_0(d\boldsymbol{\theta}). \end{aligned} \quad (279)$$

By our assumption, the RHS is finite for $t_0 > 0$. Hence, by taking t_0 large enough, the value of $\int_{t_0}^\infty \int_{t_0}^t \|\Delta \Gamma_{t,s}\|_{\hat{\nu}}^2 ds dt$ can be made arbitrarily close to zero.

The second goal is to show that $\forall t_0 > 0$,

$$\lim_{T \rightarrow \infty} \int_{t_0}^T \int_{t_0}^t \|\Gamma_{t,s}\|^2 ds dt = 0. \quad (280)$$

As a first step, we show that

$$\lim_{T \rightarrow \infty} \int_{t_0}^T \int_0^{t_0} \|\Gamma_{t-s}^\infty\|_{\hat{\nu}}^2 ds dt = 0 \quad (281)$$

because $\forall \eta \in \mathcal{W}_L(\Omega)$, there is

$$\begin{aligned} |\langle \eta, \Gamma_{t-s}^\infty \eta \rangle_{\hat{\nu}}| &= \int_D \langle \mathbf{b}(\boldsymbol{\theta}), e^{-t \nabla \nabla V_\infty(\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}))} \mathbf{b}(\boldsymbol{\theta}) \rangle \mu_0(d\boldsymbol{\theta}) \\ &\leq e^{-\lambda(t-s)} \int_D |\mathbf{b}(\boldsymbol{\theta})|^2 \mu_0(d\boldsymbol{\theta}) \\ &\leq e^{-\lambda(t-s)} \|\mathcal{M}_\infty\|_{\hat{\nu}} \|\eta\|_{\hat{\nu}}^2, \end{aligned} \quad (282)$$

where

$$\mathbf{b}(\boldsymbol{\theta}) = \int_\Omega \nabla \varphi(\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}), \mathbf{x}) \eta(\mathbf{x}) \hat{\nu}(d\mathbf{x}). \quad (283)$$

and $\mathcal{M}_\infty$ is defined as $\mathcal{M}_\infty := \mathcal{B}_\infty^\top \mathcal{B}_\infty$, or concretely, for $\eta \in \mathcal{W}_L(\omega)$,

$$\begin{aligned} (\mathcal{M}_\infty \eta)(\mathbf{x}) &:= \int_\Omega \left(\int_D \nabla \varphi(\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}'), \mathbf{x})^\top \nabla \varphi(\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}'), \mathbf{x}') \mu_0(d\boldsymbol{\theta}') \right) \eta(\mathbf{x}') \hat{\nu}(d\mathbf{x}') \\ &= \int_\Omega M(\mathbf{x}, \mathbf{x}', \mu_\infty) \eta(\mathbf{x}') \hat{\nu}(d\mathbf{x}'), \end{aligned} \quad (284)$$

where

$$M(\mathbf{x}, \mathbf{x}', \mu_\infty) := \int_D \nabla \varphi(\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}'), \mathbf{x}) \cdot \nabla \varphi(\boldsymbol{\Theta}_\infty(\boldsymbol{\theta}'), \mathbf{x}') \mu_0(d\boldsymbol{\theta}'). \quad (285)$$

In the ERM setting, $\mathcal{M}_\infty$ is effectively an $L \times L$ matrix. Thus,

$$\begin{aligned} \int_{t_0}^T \int_0^{t_0} \|\Gamma_{t-s}^\infty\|_{\hat{\nu}}^2 ds dt &\leq \int_{t_0}^T \int_0^{t_0} e^{-2\lambda(t-s)} \|\mathcal{M}_\infty\|_{\hat{\nu}}^2 ds dt \\ &\leq \|\mathcal{M}_\infty\|_{\hat{\nu}}^2 \int_{t_0}^T e^{-2\lambda(t-t_0)} dt \rightarrow 0 \quad \text{as } T \rightarrow \infty \end{aligned} \quad (286)$$

Hence, it is sufficient to show that

$$\lim_{T \rightarrow \infty} \int_{t_0}^T \int_0^{t_0} \|\Delta \Gamma_{t,s}\|^2 ds dt = 0. \quad (287)$$

We have

$$\begin{aligned} &\int_{t_0}^T \int_0^{t_0} \|\Delta \Gamma_{t,s}\|^2 ds dt \\ &\leq 4C|D'| \int_D \left(\int_{t_0}^T \left(\int_0^{t_0} e^{-2\lambda(t-s)} ds \right) |\Delta \boldsymbol{\Theta}_t(\boldsymbol{\theta})| dt \right. \\ &\quad \left. + \int_0^{t_0} \left(\int_{t_0}^T e^{-2\lambda(t-s)} dt \right) \left(|\Delta \boldsymbol{\Theta}_s(\boldsymbol{\theta})| + U_s(\boldsymbol{\theta})^2 \right) e^{2C'(U_s(\boldsymbol{\theta}) + \bar{U}_s)} ds \right) \mu_0(d\boldsymbol{\theta}) \\ &\leq 2C|D'| \lambda^{-1} \int_D \left(\int_{t_0}^T e^{-2\lambda(t-t_0)} |\Delta \boldsymbol{\Theta}_t(\boldsymbol{\theta})| dt \right. \\ &\quad \left. + \int_0^{t_0} e^{-2\lambda(t_0-s)} \left(|\Delta \boldsymbol{\Theta}_s(\boldsymbol{\theta})| + U_s(\boldsymbol{\theta})^2 \right) e^{2C'(U_s(\boldsymbol{\theta}) + \bar{U}_s)} ds \right) \mu_0(d\boldsymbol{\theta}) \\ &\leq 4C|D'| \lambda^{-1} \int_D \int_0^\infty \left(|\Delta \boldsymbol{\Theta}_s(\boldsymbol{\theta})| + U_s(\boldsymbol{\theta})^2 \right) e^{2C'(U_s(\boldsymbol{\theta}) + \bar{U}_s)} ds \mu_0(d\boldsymbol{\theta}) \\ &< \infty \end{aligned} \quad (288)$$

by assumption (237). Therefore,

$$\int_{t_0}^T \int_0^{t_0} \|\Delta \Gamma_{t,s}\|_{\mathcal{V}}^2 ds dt = \frac{1}{T-t_0} \int_{t_0}^T \int_0^{t_0} \|\Delta \Gamma_{t,s}\|_{\mathcal{V}}^2 ds dt \xrightarrow{T \rightarrow \infty} 0. \quad (289)$$

This concludes the proof of Lemma D.8. $\square$

D.4.2 Interpretation of the Assumption (237)

Below, we will illustrate the assumption (237)

$$Q := \int_D \int_0^\infty (|\Delta \Theta_t(\theta)| + U_t(\theta)^2) e^{C_1(U_t(\theta) + \bar{U}_t)} dt \mu_0(d\theta) < \infty, \quad (290)$$

in Theorem 3.7 by giving examples that satisfy this condition.

First, consider an example where $\exists \kappa > 0, \alpha > 1$ such that $\forall \theta \in \text{supp } \mu_0$ and $\forall t > 0$,

$$|\Delta \Theta_t(\theta)| < \kappa(t+1)^{-\alpha}, \quad (291)$$

that is, all characteristic flows share a uniform asymptotic convergence rate on the order of $t^{-\alpha}$. Then $\forall \theta \in \text{supp } \mu_0$,

$$U_t(\theta) = \int_t^\infty |\Delta \Theta_s(\theta)| ds \leq \frac{\kappa}{\alpha-1} (t+1)^{-(\alpha-1)} \quad (292)$$

and thus

$$\bar{U}_t \leq \frac{\kappa}{\alpha-1} (t+1)^{-(\alpha-1)}. \quad (293)$$

Therefore,

$$\begin{aligned} Q &\leq \int_D \int_0^\infty (|\Delta \Theta_t(\theta)| + U_t(\theta)^2) e^{C_1(U_t(\theta) + \bar{U}_t)} dt \mu_0(d\theta) \\ &\leq \int_0^\infty \left(\kappa(t+1)^{-\alpha} + \left(\frac{\kappa}{\alpha-1} \right)^2 (t+1)^{-2(\alpha-1)} \right) e^{\frac{2C_1\kappa}{\alpha-1}} dt, \end{aligned} \quad (294)$$

which is finite as long as $\alpha > \frac{3}{2}$. Thus,

Proposition D.10. *If $\exists \kappa > 0, \alpha > \frac{3}{2}$ such that $\forall \theta \in \text{supp } \mu_0$ and $\forall t \geq 0$,*

$$|\Delta \Theta_t(\theta)| = |\Theta_t(\theta) - \Theta_\infty(\theta)| < \kappa(t+1)^{-\alpha}, \quad (295)$$

then the condition (237) is satisfied.

Moreover, the assumption allows flexibility in having non-uniform convergence rate for different characteristic flows, $\Theta_t(\theta)$. Suppose that $\exists \kappa : \text{supp } \mu_0 \rightarrow \mathbb{R}_+$ and $\alpha > \frac{3}{2}$ such that $\forall \theta \in \text{supp } \mu_0$,

$$|\Delta \Theta_t(\theta)| < \kappa(\theta)(t+1)^{-\alpha}. \quad (296)$$

Then

$$U_t(\theta) = \int_t^\infty |\Delta \Theta_s(\theta)| ds \leq \frac{\kappa}{\alpha-1} (t+1)^{-(\alpha-1)} \quad (297)$$

and so

$$\begin{aligned}
Q &\leq \int_D \int_0^\infty (|\Delta \Theta_t(\boldsymbol{\theta})| + U_t(\boldsymbol{\theta})^2) e^{2C_1(U_0(\boldsymbol{\theta}))} dt \mu_0(d\boldsymbol{\theta}) \\
&\leq \int_D \int_0^\infty \left(\kappa(\boldsymbol{\theta})(t+1)^{-\alpha} + \left(\frac{\kappa(\boldsymbol{\theta})}{\alpha-1} \right)^2 (t+1)^{-2(\alpha-1)} \right) e^{\frac{2C_1\kappa(\boldsymbol{\theta})}{\alpha-1}} dt \\
&\leq C_2 \int_D (\kappa(\boldsymbol{\theta}) + \kappa(\boldsymbol{\theta})^2) e^{\frac{2C_1\kappa(\boldsymbol{\theta})}{\alpha-1}} \mu_0(d\boldsymbol{\theta}) .
\end{aligned} \tag{298}$$

Therefore,

Proposition D.11. Suppose $\exists \alpha > \frac{3}{2}$ and a function $\kappa : \text{supp } \mu_0 \rightarrow \mathbb{R}_+$, which satisfies

$$\int_D \left(\kappa(\boldsymbol{\theta}) + \kappa(\boldsymbol{\theta})^2 \right) e^{\frac{2C_1\kappa(\boldsymbol{\theta})}{\alpha-1}} \mu_0(d\boldsymbol{\theta}) < \infty, \tag{299}$$

such that $\forall \boldsymbol{\theta} \in \text{supp } \mu_0$,

$$|\Delta \Theta_t(\boldsymbol{\theta})| = |\Theta_t(\boldsymbol{\theta}) - \Theta_\infty(\boldsymbol{\theta})| \leq \kappa(\boldsymbol{\theta})(t+1)^{-\alpha}. \tag{300}$$

Then the condition (237) is satisfied.

D.4.3 Relationship between Theorem 3.7 and [11]

As a comparison to our result, Chizat [11, Theorem 3.8] shows that under assumptions including (51) as well as the uniqueness and sparseness of the global minimizer, an alternative type of particle gradient descent (with a different homogeneity degree in the loss function and under the conic metric, which give rise to gradient flow in Wasserstein-Fisher-Rao metric instead of Wasserstein metric) converges to the global minimizer for large enough n (depending exponentially on d) with a uniform rate. This implies that in that setting, $\lim_{t \rightarrow \infty} \lim_{n \rightarrow \infty} n \|f_t^{(n)} - f_t\|_{\mathcal{V}}^2 = \lim_{n \rightarrow \infty} \lim_{t \rightarrow \infty} n \|f_t^{(n)} - f_t\|_{\mathcal{V}}^2 = 0$, $\mathbb{P}_0$ -almost surely.

E Properties of the Minimizers of the Regularized Loss

First, under Assumption 2.1, i.e. in the shallow neural networks setting, define

$$\hat{F}(\mathbf{z}) = \int_{\Omega} f_*(\mathbf{x}) \hat{\varphi}(\mathbf{z}, \mathbf{x}) \hat{\nu}(d\mathbf{x}), \quad \hat{K}(\mathbf{z}, \mathbf{z}') = \int_{\Omega} \hat{\varphi}(\mathbf{z}, \mathbf{x}) \hat{\varphi}(\mathbf{z}', \mathbf{x}) \hat{\nu}(d\mathbf{x}). \tag{301}$$

and

$$\hat{V}(\mathbf{z}, \mu) = -\hat{F}(\mathbf{z}) + \int_D c' \hat{K}(\mathbf{z}, \mathbf{z}') \mu(dc', d\mathbf{z}'). \tag{302}$$

We prove Proposition 3.9, which we extend into:

Proposition E.1. Under Assumptions 2.1, 2.2, and 3.8, the minimizers of the loss $\mathcal{L}(\mu)$ defined in (4) are all in the form

$$\mu_\lambda(dc, d\mathbf{z}) = \delta_{c_\lambda}(dc) \hat{\mu}_+(d\mathbf{z}) + \delta_{-c_\lambda}(dc) \hat{\mu}_-(d\mathbf{z}) \tag{303}$$

where $c_\lambda \geq 0$ and $\hat{\mu}_\pm \in \mathcal{P}(\hat{D})$ satisfy

$$\begin{aligned} \forall \mathbf{z} \in \text{supp } \hat{\mu}_- & : -\hat{F}(\mathbf{z}) + c_\lambda \int_{\hat{D}} \hat{K}(\mathbf{z}, \mathbf{z}') (\hat{\mu}_+(d\mathbf{z}') - \hat{\mu}_-(d\mathbf{z}')) = \lambda c_\lambda, \\ \forall \mathbf{z} \in \text{supp } \hat{\mu}_+ & : -\hat{F}(\mathbf{z}) + c_\lambda \int_{\hat{D}} \hat{K}(\mathbf{z}, \mathbf{z}') (\hat{\mu}_+(d\mathbf{z}') - \hat{\mu}_-(d\mathbf{z}')) = -\lambda c_\lambda, \\ \forall \mathbf{z} \in \hat{D} & : \left| -\hat{F}(\mathbf{z}) + c_\lambda \int_{\hat{D}} \hat{K}(\mathbf{z}, \mathbf{z}') (\hat{\mu}_+(d\mathbf{z}') - \hat{\mu}_-(d\mathbf{z}')) \right| \leq \lambda c_\lambda. \end{aligned} \quad (304)$$

In addition, the constant c_λ is unique and positive if $\hat{F}(\mathbf{z})$ is not identically zero on $\hat{D}$, the closure of the supports of $\hat{\mu}_\pm$ are disjoint (i.e. $\overline{\text{supp } \hat{\mu}_+} \cap \overline{\text{supp } \hat{\mu}_-} = \emptyset$), and the function

$$f_\lambda = \int_D c \hat{\varphi}(\mathbf{z}, \cdot) \mu_\lambda(dc, d\mathbf{z}) = c_\lambda \int_{\hat{D}} \hat{\varphi}(\mathbf{z}, \cdot) (\hat{\mu}_+(d\mathbf{z}) - \hat{\mu}_-(d\mathbf{z})) \quad (305)$$

is the same for all minimizers and satisfies

$$\frac{1}{4} \lambda^2 |c_\lambda|^2 \hat{K}_M^{-1} \leq \|f_* - f_\lambda\|_{\hat{V}}^2, \quad \|f_* - f_\lambda\|_{\hat{V}}^2 + \lambda |c_\lambda|^2 \leq \lambda |\gamma_1(f_*)|^2. \quad (306)$$

where $\hat{K}_M = \max_{\mathbf{z} \in \hat{D}} \|\hat{\varphi}(\mathbf{z}, \cdot)\|_{\hat{V}}^2 = \max_{\mathbf{z} \in \hat{D}} \hat{K}(\mathbf{z}, \mathbf{z})$.

Remark E.2. Note that the proposition automatically implies that $\gamma_1(f_\lambda) \leq \gamma_1(f_*) < \infty$. It also implies that

$$\int_D |c|^q \mu_\lambda(dc, d\mathbf{z}) = |c_\lambda|^q = |\gamma_\lambda|_{TV}^q \leq |\gamma_1(f_*)|^q \quad \forall q \in \mathbb{R}_+ \quad (307)$$

where $\gamma_\lambda = \int_{\mathbb{R}} c \mu_\lambda(dc, \cdot)$. Finally note that the proposition holds if we replace the empirical loss by the population loss.

Proof: The fact that this loss can only be minimized by minimizers follows from the compactness of the sets $\{\mu \in \mathcal{P}(D) : \mathcal{L}(\mu) \leq u, u \in \mathbb{R}\}$. The minimizers of $\mathcal{L}(\mu)$ must satisfy the following Euler-Lagrange equations [57]:

$$\forall (c, \mathbf{z}) \in D : -c \hat{F}(\mathbf{z}) + c \int_D c' \hat{K}(\mathbf{z}, \mathbf{z}') \mu(dc', d\mathbf{z}') + \frac{1}{2} \lambda |c|^2 \equiv c \hat{V}(\mathbf{z}) + \frac{1}{2} \lambda |c|^2 \geq \bar{V}, \quad (308)$$

with equality on the support of μ and where $\bar{V}$ is the expectation of the left hand side with respect to $\mu(dc, d\mathbf{z})$. Minimizing the left hand side of (308) over c at fixed $\mathbf{z}$, we deduce that

$$\forall \mathbf{z} \in \hat{D} : \min_c \left(c \hat{V}(\mathbf{z}) + \frac{1}{2} \lambda |c|^2 \right) \geq \bar{V}, \quad (309)$$

with equality for $\mathbf{z}$ in the support of $\hat{\mu} = \int_{\mathbb{R}} \mu(dc, \cdot)$. This means that for any $\mathbf{z} \in \text{supp } \hat{\mu}$, there can only be one $c = c(\mathbf{z})$ in $\text{supp } \mu$, with $c(\mathbf{z})$ satisfying the Euler-Lagrange equation associated with (309)

$$\hat{V}(\mathbf{z}) + \lambda c(\mathbf{z}) = 0 \quad \Leftrightarrow \quad \hat{V}(\mathbf{z}) = -\lambda c(\mathbf{z}) \quad (310)$$

If we insert this equality back in $c(\mathbf{z}) \hat{V}(\mathbf{z}) + \frac{1}{2} \lambda |c(\mathbf{z})|^2 = \bar{V}$, we deduce that $|c(\mathbf{z})| = c_\lambda$, with the constant c_λ related to $\bar{V}$ as

$$-\frac{1}{2} \lambda |c_\lambda|^2 = \bar{V}, \quad (311)$$

and furthermore, $\forall \mathbf{z} \in \text{supp } \hat{\mu}$,

$$\hat{V}(\mathbf{z}) = \begin{cases} -\lambda c_\lambda & \text{if } c(\mathbf{z}) = c_\lambda \\ \lambda c_\lambda & \text{if } c(\mathbf{z}) = -c_\lambda \end{cases}. \quad (312)$$

These considerations imply that the minimizer must be of the form (303), and if we combine (309) and (311) and evaluate the minimum on c explicitly we deduce that $\hat{\mu}_\pm$ and c_λ must satisfy the equations in (304). It is also clear from (304) that we must have $\overline{\text{supp } \hat{\mu}_+} \cap \overline{\text{supp } \hat{\mu}_-} = \emptyset$: indeed if there was a point $z \in \overline{\text{supp } \hat{\mu}_+} \cap \overline{\text{supp } \hat{\mu}_-}$, then at that point $\hat{V}(z)$ would be discontinuous, which is not possible since this function is continuously differentiable for any μ by our assumptions on $\hat{\varphi}$. Finally, to show that we must have that $c_\lambda > 0$ if $F(z)$ is not identically zero on $\hat{D}$, note that if $c_\lambda = 0$, (308) reduces to

$$\forall (c, z) \in D \quad : \quad -c\hat{F}(z) + \frac{1}{2}\lambda|c|^2 \geq 0 \quad (313)$$

which can only be satisfied if $\hat{F}(z) = 0$.

To show that c_λ and the function in (305) are unique, let μ_λ and μ'_λ be two different minimizers and consider

$$f_\lambda = \int_D c\hat{\varphi}(z, \cdot)\mu_\lambda(dc, dz) \quad \text{and} \quad f'_\lambda = \int_D c\hat{\varphi}(z, \cdot)\mu'_\lambda(dc, dz) \quad (314)$$

Let us evaluate the loss on $a\mu_\lambda + (1-a)\mu'_\lambda \in \mathcal{P}(D)$ with $a \in [0, 1]$. By convexity of $\mathcal{E}_\lambda$ we have

$$\mathcal{L}(a\mu_\lambda + (1-a)\mu'_\lambda) \leq a\mathcal{L}(\mu_\lambda) + (1-a)\mathcal{L}(\mu'_\lambda) = \mathcal{L}(\mu_\lambda) = \mathcal{L}(\mu'_\lambda) \quad (315)$$

Since $a\mu_\lambda + (1-a)\mu'_\lambda$ cannot have a lower loss than this minimum, we must have equality in (315), which reduces to

$$\begin{aligned} & \|f_* - af_\lambda - (1-a)f'_\lambda\|_{\hat{\nu}}^2 + a\lambda|c_\lambda|^2 + (1-a)\lambda|c'_\lambda|^2 \\ &= \|f_* - f_\lambda\|_{\hat{\nu}}^2 + \lambda|c_\lambda|^2 \\ &= \|f_* - f'_\lambda\|_{\hat{\nu}}^2 + \lambda|c'_\lambda|^2, \end{aligned} \quad (316)$$

where c_λ and c'_λ are associated with μ_λ and μ'_λ , respectively. Clearly these equations can only be fulfilled for all $a \in [0, 1]$ if $c_\lambda = c'_\lambda$ and $f_\lambda = f'_\lambda$ $\hat{\nu}$ -a.e. on Ω .

To establish (306), notice that if μ_λ is a minimizer and f_λ is given by (305), then we can derive from (312) that

$$-\int_\Omega f_\lambda(x)f_*(x)\hat{\nu}(dx) + \|f_\lambda\|_{\hat{\nu}}^2 + \lambda|c_\lambda|^2 = 0. \quad (317)$$

This gives, using Cauchy-Schwartz,

$$\lambda|c_\lambda|^2 = \int_\Omega f_\lambda(x)(f_*(x) - f_\lambda(x))\hat{\nu}(dx) \leq \|f_\lambda\|_{\hat{\nu}} \|f_* - f_\lambda\|_{\hat{\nu}}. \quad (318)$$

Now notice that

$$\|f_\lambda\|_{\hat{\nu}}^2 = c_\lambda^2 \int_{\hat{D} \times \hat{D}} \hat{K}(z, z') (\hat{\mu}_+(dz) - \hat{\mu}_-(dz)) (\hat{\mu}_+(dz') - \hat{\mu}_-(dz')) \leq 4c_\lambda^2 \hat{K}_M. \quad (319)$$

Using (319) in (318) and reorganizing gives the first inequality in (306). To establish the second, let $\mu_* \in \mathcal{M}_+(D)$ be the measure that minimizes $\int_D |c|\mu(dc, dz)$ under the constraint that $f_* = \int_D c\hat{\varphi}(z, \cdot)\mu_*(dc, dz)$, so that $\int_D |c|\mu_*(dc, dz) = \gamma_1(f_*)$ —the measure μ_* exists since we assumed that $f_* \in \mathcal{F}_1$. Evaluated on μ_* , the loss is

$$\mathcal{L}(\mu_*) = \lambda|\gamma_1(f_*)|^2. \quad (320)$$

Any minimizer μ_λ of $\mathcal{L}(\mu)$ must do at least as well, i.e we must have

$$\|f_* - f_\lambda\|_{\hat{\nu}}^2 + \lambda \int_D |c|^2 \mu_\lambda(dc, dz) = \|f_* - f_\lambda\|_{\hat{\nu}}^2 + \lambda|c_\lambda|^2 \leq \lambda|\gamma(f_*)|^2. \quad (321)$$

This establishes the second inequality in (306). $\square$

F Analytical Calculations of the Resampling Error

Derivations similar to the one presented here can be found in [5, 15, 53]. In the setting of ReLU without bias on unit sphere, we take $\hat{D} = \Omega = \mathbb{S}^d \subseteq \mathbb{R}^{d+1}$, $\hat{\varphi}(\mathbf{z}, \mathbf{x}) = \max(\langle \mathbf{z}, \mathbf{x} \rangle, 0)$, and ν is equal to the uniform measure on $\mathbb{S}^d$. In this case,

$$\hat{K}(\mathbf{z}, \mathbf{z}') = \int_{\Omega} \hat{\varphi}(\mathbf{z}, \mathbf{x}) \hat{\varphi}(\mathbf{z}', \mathbf{x}) \nu(d\mathbf{x}) = \frac{1}{2(d+1)\pi} (\sin \alpha + (\pi - \alpha) \cos \alpha), \quad (322)$$

with α being the angle between $\mathbf{z}$ and $\mathbf{z}'$, and

$$\int_{\Omega} |\hat{\varphi}(\mathbf{z}, \mathbf{x})|^2 \nu(d\mathbf{x}) = \frac{1}{2} \int_{\Omega} (\langle \mathbf{x}, \mathbf{z} \rangle)^2 \nu(d\mathbf{x}) = \frac{1}{2(d+1)} \quad (323)$$

Thus, taking μ_* to be the measure representing the teacher network, $\mu_* = \frac{1}{m_t} \sum_{i=1}^{m_t} \delta_{\mathbf{z}_i}(d\mathbf{z}) \delta_1(dc)$, we have

$$\begin{aligned} \int_D \|\varphi(\boldsymbol{\theta}, \cdot)\|_{\nu}^2 \mu_*(d\boldsymbol{\theta}) &= \int_D \int_{\Omega} |\varphi(\boldsymbol{\theta}, \mathbf{x})|^2 \nu(d\mathbf{x}) \mu_*(d\boldsymbol{\theta}) \\ &= \int_D \frac{c^2}{2(d+1)} \mu_*(d\boldsymbol{\theta}) \\ &= \frac{1}{2(d+1)} \end{aligned} \quad (324)$$

On the other hand,

$$\begin{aligned} \|f_*\|_{\nu}^2 &= \int_{\Omega} \left| \int_D \varphi(\boldsymbol{\theta}, \mathbf{x}) \mu_*(d\boldsymbol{\theta}) \right|^2 \nu(d\mathbf{x}) \\ &= \int_D \int_D cc' \hat{K}(\mathbf{z}, \mathbf{z}') \mu_*(d\boldsymbol{\theta}) \mu_*(d\boldsymbol{\theta}') \\ &= \frac{1}{m_t^2} \sum_{i,j=1}^{m_t} \hat{K}(\mathbf{z}_i, \mathbf{z}_j) \end{aligned} \quad (325)$$

In the experiments described in the main text, we take $m_t = 2$, and $\mathbf{z}_1$ and $\mathbf{z}_2$ are initialized with a fixed random seed such that their angle, α_{12} , equal to 1.766. Thus,

$$\|f_*\|_{\nu}^2 = \frac{1}{4(d+1)\pi} (0 + \pi) + \frac{1}{4(d+1)\pi} (\sin \alpha_{12} + (\pi - \alpha_{12}) \cos \alpha_{12}) \approx 0.012 \quad (326)$$

Together, we get a numerical value of the RHS of (46) if we replace μ_{∞} , f_{∞} and $\hat{\nu}$ by μ_* , f_* and ν , respectively.