
A Solvable High-Dimensional Model of GAN

Chuang Wang^{1,2}
wangchuang@ia.ac.cn

Hong Hu²
honghu@g.harvard.edu

Yue M. Lu²
yuelu@seas.harvard.edu

1. State Key Laboratory of Pattern Recognition, Institute of Automation,
Chinese Academy of Science, 95 Zhong Guan Cun Dong Lu, Beijing 100190, China
2. John A. Paulson School of Engineering and Applied Sciences, Harvard University
33 Oxford Street, Cambridge, MA 02138, USA

Abstract

We present a theoretical analysis of the training process for a single-layer GAN fed by high-dimensional input data. The training dynamics of the proposed model at both microscopic and macroscopic scales can be exactly analyzed in the high-dimensional limit. In particular, we prove that the macroscopic quantities measuring the quality of the training process converge to a deterministic process characterized by an ordinary differential equation (ODE), whereas the microscopic states containing all the detailed weights remain stochastic, whose dynamics can be described by a stochastic differential equation (SDE). This analysis provides a new perspective different from recent analyses in the limit of small learning rate, where the microscopic state is always considered deterministic, and the contribution of noise is ignored. From our analysis, we show that the level of the background noise is essential to the convergence of the training process: setting the noise level too strong leads to failure of feature recovery, whereas setting the noise too weak causes oscillation. Although this work focuses on a simple copy model of GAN, we believe the analysis methods and insights developed here would prove useful in the theoretical understanding of other variants of GANs with more advanced training algorithms.

1 Introduction

A generative adversarial network (GAN) [1] seeks to learn a high-dimensional probability distribution from samples. While there have been numerous advances on the application front [2–6], considerably less is known about the underlying theory and conditions that can explain or guarantee the successful trainings of GANs.

Recently, it has been a very active area of research to study either the equilibrium properties [7–9] or the training dynamics [10, 11]. Specifically, there is a line of works studying the dynamics of the gradient-based training algorithms *e.g.*, [11–16]. The basic idea is the following. The evolution of the learnable parameters in the training dynamics can be considered as a discrete-time process. With a proper time scaling, this discrete-time process converges to a deterministic continuous-time process as the learning rates tend to 0, which is characterized by an ordinary differential equation (ODE). By studying local stability of the ODE’s fixed points, [12] shows that oscillation in the training algorithm is due to the eigenvalues of the Jacobian of the gradient vector field with zero real part and large imaginary part. Due to this fact, various stabilization approaches are proposed, for example adding additional regularizers [13, 14], and using two timescale [15] training. Very recently, [16] argues that those stabilization techniques may encourage the algorithms to converge non-Nash stationary points. All above works consider a small-learning-rates limit, where the limiting process

is always deterministic. The stochasticity and the effect of the noise is essentially ignored, which may not reflect practical situations. Thus, a new analysis paradigm to study the dynamics with the consideration of the intrinsic stochasticity is needed.

In this paper, we present a *high-dimensional* and *exactly solvable* model of GAN. Its dynamics can be precisely characterized at both macroscopic and microscopic scales, where the former is deterministic and the latter remains stochastic. Interestingly, our theoretical analysis shows that injecting additional noise can stabilize the training. Specifically, our main technical contributions are twofold:

- We present an asymptotically exact analysis of the training process of the proposed GAN model. Our analysis is carried out on both the *macroscopic* and the *microscopic* levels. The macroscopic state measures the overall performance of the training process, whereas the microscopic state contains all the detailed weights information. In the high-dimensional limit ($n \rightarrow \infty$), we show that the former converges to a deterministic process governed by an ordinary differential equation (ODE), whereas the latter stays stochastic described by a stochastic differential equation (SDE).
- We show that depending on the choice of the learning rates and the strength of noise, the training process can reach either a successful, a failed, an oscillating, or a mode-collapsing phase. By studying the stabilities of the fixed points of the limiting ODEs, we precisely characterize when each phase takes place. The analysis reveals a condition on the learning rates and the noise strength for successful training. We show that the level of the background noise is essential to the convergence of the training process: setting the noise level too strong (small signal-to-noise ratio) leads to failure of feature recovery, whereas setting the noise too weak (large signal-to-noise ratio) causes oscillation.

Our work builds upon a general analysis framework [17] for studying the scaling limits of high-dimensional exchangeable stochastic processes with applications to nonlinear regression problems. Similar techniques have also been used in the literature to study Monte Carlo methods [18], online perceptron learning [19, 20], online sparse PCA [21], subspace estimation [22], online ICA [23] and more recently, the supervised learning of two-layer neural networks [24], but to our best knowledge, this technique has not yet been used in analyzing GANs.

The rest of the paper is organized as follows. We present the proposed GAN model and the associated training algorithm in Section 2. Our main results are presented in Section 3, where we show that the macroscopic and microscopic dynamics of the training process converge to their respective limiting processes that are characterized by an ODE and SDE, respectively. In Section 4, we analyze the stationary solutions of the limiting ODEs and precisely characterizes the long-term behaviors of the training process. We conclude in Section 5.

2 Formulations

In this section, we introduce the proposed GAN model and specify the associated training algorithm.

Model for the real data. In order to establish the theoretical analysis, we first impose a model for the probability distribution from which we draw our real data samples. We assume that the real data $\mathbf{y}_k \in \mathbb{R}^n$, $k = 0, 1, \dots$ are drawn according to the following generative model:

$$\mathbf{y}_k = \mathcal{G}(\mathbf{c}_k, \mathbf{a}_k; \mathbf{U}, \eta_T) \stackrel{\text{def}}{=} \mathbf{U}\mathbf{c}_k + \sqrt{\eta_T}\mathbf{a}_k, \quad (1)$$

where $\mathbf{U} \in \mathbb{R}^{n \times d}$ is a deterministic unknown feature matrix with d features; $\mathbf{c}_k \in \mathbb{R}^d$ is a random vector drawn from an unknown distribution $\mathcal{P}_c$; $\mathbf{a}_k$ is an n -dimensional random vector acting as the background noise; and η_T is a parameter to control the strength of noise. Without loss of generality¹, we assume $\mathbf{U}^\top \mathbf{U} = \mathbf{I}_d$, where $\mathbf{I}_d$ is the $d \times d$ identity matrix.

This generative model, referred to as the spiked covariance model [25] in the literature, is commonly used in the theoretical study of principal component analysis (PCA). We note that this model is not a trivial task for PCA even when $d = 1$ if the variance of the noise $\mathbf{a}_k$ is a non-zero constant. As

¹If $\mathbf{U}$ is not orthogonal, we can rewrite $\mathbf{U}\mathbf{c}$ in (1) as $(\mathbf{U}\mathbf{R})(\mathbf{R}^{-1}\mathbf{c})$, where $\mathbf{R}$ is a matrix that orthogonalizes and normalizes the columns of $\mathbf{U}$. We can then study an equivalent system where the new feature vector is $\mathbf{R}^{-1}\mathbf{c}$.

proved in [25], the best estimator can not perfectly recover the signal $\mathbf{U}$ given an $\mathcal{O}(n)$ number of samples $\mathbf{y}_k$. Thus, it is of sufficient interest to investigate whether a GAN can retrieve informative results for the principal components in the same scaling limit.

The GAN model The GAN we are going to analyze is defined as follows. We assume that the generator $\mathcal{G}$ has the same linear structure as the real data model (1) given above:

$$\tilde{\mathbf{y}}_k = \mathcal{G}(\tilde{\mathbf{c}}_k, \tilde{\mathbf{a}}_k; \mathbf{V}, \eta_G) \quad (2)$$

but the parameters are different. Here, $\tilde{\mathbf{y}}_k$ denotes a fake sample produced by the generator; $\tilde{\mathbf{a}}_k$ is an n -dimensional random noise vector; the random variable $\tilde{\mathbf{c}}_k$ is drawn from a fixed distribution $\mathcal{P}_{\tilde{\mathbf{c}}}$; η_G is the noise strength; and the matrix $\mathbf{V} \in \mathbb{R}^{n \times d}$ represents the parameters of the generator. (In an ideal case in which the generator learns the underlying true probability distribution perfectly, we have $\mathbf{V} = \mathbf{U}$.) Throughout the paper, we follow the notational convention that all the symbols that are decorated with a tilde (e.g., $\tilde{\mathbf{y}}_k, \tilde{\mathbf{c}}_k, \tilde{\mathbf{a}}_k$) denote quantities associated with the generator.

We define the discriminator $\mathcal{D}$ of our GAN model as

$$\mathcal{D}(\mathbf{y}; \mathbf{w}) \stackrel{\text{def}}{=} \hat{D}(\mathbf{y}^\top \mathbf{w}).$$

Here, $\mathbf{y}$ is an input vector, which can be either the real data $\mathbf{y}_k$ from (1) or the fake one $\tilde{\mathbf{y}}_k$ from (2); $\hat{D} : \mathbb{R} \mapsto \mathbb{R}$ can be any function; and the vector $\mathbf{w} \in \mathbb{R}^n$ represents the parameters associated with the discriminator. Later, we will show that the generator can learn multiple features even though the discriminator only has one feature vector $\mathbf{w}$. Discriminators with multiple features can also be analyzed in a similar way, but in this paper we consider the single-feature discriminator for simplicity.

The training algorithm. The proposed GAN model has two set of parameters $\mathbf{V}$ and $\mathbf{w}$ to be learned from the data. The training process is formulated as the following MinMax problem

$$\min_{\mathbf{V}} \max_{\mathbf{w}} \mathbb{E}_{\mathbf{y} \sim \mathcal{P}(\mathbf{y}; \mathbf{U})} \mathbb{E}_{\tilde{\mathbf{y}} \sim \tilde{\mathcal{P}}(\tilde{\mathbf{y}}; \mathbf{V})} \mathcal{L}(\mathbf{y}, \tilde{\mathbf{y}}; \mathbf{w}), \quad (3)$$

where the two probability distributions $\mathcal{P}(\mathbf{y}; \mathbf{U})$ and $\tilde{\mathcal{P}}(\tilde{\mathbf{y}}; \mathbf{V})$ represent the distributions of the real data $\mathbf{y}$ and the fake data $\tilde{\mathbf{y}}$ as specified by (1) and (2) respectively, and

$$\mathcal{L}(\mathbf{y}, \tilde{\mathbf{y}}; \mathbf{w}) \stackrel{\text{def}}{=} F(\hat{D}(\mathbf{y}^\top \mathbf{w})) - \tilde{F}(\hat{D}(\tilde{\mathbf{y}}^\top \mathbf{w})) - \frac{\lambda}{2} H(\mathbf{w}^\top \mathbf{w}) + \frac{\lambda}{2} \text{tr}(H(\mathbf{V}^\top \mathbf{V})) \quad (4)$$

with $F(\cdot)$ and $\tilde{F}(\cdot)$ being two functions that quantify the performance of the discriminator and $\lambda > 0$ being a constant. The function $H(\cdot)$ acts as a regularization term introduced to control the magnitude of the parameters $\mathbf{w}$ and $\mathbf{V}$. It can be an arbitrary real-valued function, which is applied element-wisely if the input is a matrix.

We consider a standard training algorithm that uses the vanilla stochastic gradient descent/ascent (SGDA) to seek a solution of (3). To simplify the theoretical analysis, we consider an online (i.e., streaming) setting where each data sample $\mathbf{y}_k$ is used only once. At step k , the model parameters $\mathbf{w}_k$ and $\mathbf{V}_k$ are updated using a new real sample $\mathbf{y}_k$ and two fake samples $\tilde{\mathbf{y}}_{2k}$ and $\tilde{\mathbf{y}}_{2k+1}$, according to

$$\begin{aligned} \mathbf{w}_{k+1} &= \mathbf{w}_k + \frac{\tau}{n} \nabla_{\mathbf{w}_k} \mathcal{L}(\mathbf{y}_k, \tilde{\mathbf{y}}_{2k}; \mathbf{w}_k) \\ \mathbf{V}_{k+1} &= \mathbf{V}_k - \frac{\tilde{\tau}}{n} \nabla_{\mathbf{V}_k} \mathcal{L}(\mathbf{y}_k, \mathcal{G}(\tilde{\mathbf{c}}_{2k+1}, \tilde{\mathbf{a}}_{2k+1}; \mathbf{V}_k; \eta_G); \mathbf{w}_k), \end{aligned} \quad (5)$$

where $\tilde{\mathbf{c}}_{2k+1}, \tilde{\mathbf{a}}_{2k+1}$ are random variables that generates the fake sample $\tilde{\mathbf{y}}_{2k+1}$ according to (2). The two parameters τ and $\tilde{\tau}$ in the above expressions control the learning rates of the discriminator and the generator, respectively. In (5), we only consider a single-step update for $\mathbf{w}_k$. This is a special case of Algorithm 1 in [1] with the batch-size m set to 1. We note that the analysis presented in this paper can be naturally extended to the mini-batch case where m is a finite number.

Example 1. We define $F(\hat{D}(x)) = \tilde{F}(\hat{D}(x)) = x^2/2$, and the regularizer function $H(\mathbf{A}) = \log \cosh(\mathbf{A} - \mathbf{I})$, where $\mathbf{I}$ is the identity matrix with the same dimension of $\mathbf{A}$, and the function $\log \cosh(\cdot)$ transforms the input matrix element-wisely. We use this specific regularizer to control the magnitude of the model parameters $\mathbf{V}$ and $\mathbf{w}$. In practice, any convex function with its minimum reached at zero would be fine. Our choice $\log \cosh(\mathbf{A} - \mathbf{I})$ here is just a convenient special case since its derivative $H'(x) = \tanh(x)$ is smooth and bounded. Furthermore, we set the regularization parameter $\lambda \rightarrow \infty$, the original problem (3) becomes a constrained MinMax problem

$$\min_{\text{diag}(\mathbf{V}^\top \mathbf{V}) = \mathbf{I}_d} \max_{\|\mathbf{w}\|=1} \mathbb{E}_{\mathbf{y} \sim \mathcal{P}} \mathbb{E}_{\tilde{\mathbf{y}} \sim \tilde{\mathcal{P}}} \left[(\mathbf{y}^\top \mathbf{w})^2 - (\tilde{\mathbf{y}}^\top \mathbf{w})^2 \right],$$

in which the diagonal operation $\text{diag}(\mathbf{A})$ returns a matrix where the diagonal entries are the same as $\mathbf{A}$ and the off-diagonal entries are all zero. The condition $\text{diag}(\mathbf{V}^\top \mathbf{V}) = \mathbf{I}_d$ ensures that each column vector of $\mathbf{V}$ is normalized.

3 Dynamics of the GAN

Definition 1. Let $\mathbf{X}_k \stackrel{\text{def}}{=} [\mathbf{U}, \mathbf{V}_k, \mathbf{w}_k] \in \mathbb{R}^{n \times (2d+1)}$. We call $\mathbf{X}_k$ the *microscopic state* of the training process at iteration step k .

The microscopic state $\mathbf{X}_k$ contains all the information about the training process. In fact, the sequence $\{\mathbf{X}_k\}_{k=0,1,2,\dots}$ forms a Markov chain on $\mathbb{R}^{n \times (2d+1)}$. This can be easily verified from the update rule of $\mathbf{X}_k$ as defined in (5), in which the real data $\mathbf{y}_k$ and fake data $\tilde{\mathbf{y}}_k$ are drawn according to (1) and (2) respectively. The Markov chain is driven by the initial state $\mathbf{X}_0$ and the sequence of random variables $\{(c_k, \mathbf{a}_k, \tilde{c}_{2k}, \tilde{\mathbf{a}}_{2k}, \tilde{c}_{2k+1}, \tilde{\mathbf{a}}_{2k+1})\}_{k=0,1,2,\dots}$.

Definition 2. Let $\mathbf{P}_k \stackrel{\text{def}}{=} \mathbf{U}^\top \mathbf{V}_k$, $\mathbf{q}_k \stackrel{\text{def}}{=} \mathbf{U}^\top \mathbf{w}_k$, $\mathbf{r}_k \stackrel{\text{def}}{=} \mathbf{V}_k^\top \mathbf{w}_k$, $\mathbf{S}_k \stackrel{\text{def}}{=} \mathbf{V}_k^\top \mathbf{V}_k$, and $z_k \stackrel{\text{def}}{=} \mathbf{w}_k^\top \mathbf{w}_k$. We call the tuple $\{\mathbf{P}_k, \mathbf{q}_k, \mathbf{r}_k, \mathbf{S}_k, z_k\}$ the *macroscopic state* of the Markov chain $\mathbf{X}_k$ at step k .

Those macroscopic quantities measure the cosine similarities among the feature vectors of the true model $\mathbf{U}$, the generator $\mathbf{V}_k$ and the discriminator $\mathbf{w}_k$. For example, the cosine of the angle between the i th true feature (*i.e.*, the i th column of $\mathbf{U}$) and the j th feature estimated in the generator (*i.e.*, the j th column of $\mathbf{V}_k$) is $[\mathbf{P}_k]_{i,j} / \sqrt{[\mathbf{S}_k]_{j,j}}$, where $[\mathbf{P}_k]_{i,j}$ is the inner product between the two feature vectors and $\sqrt{[\mathbf{S}_k]_{j,j}}$ is the norm of the j th column of $\mathbf{V}_k$. (The columns of $\mathbf{U}$ are unit vectors and need not be normalized here.) For simplicity, we introduce a compact notation for the macroscopic state:

$$\mathbf{M}_k \stackrel{\text{def}}{=} \mathbf{X}_k^\top \mathbf{X}_k = \begin{bmatrix} \mathbf{I} & \mathbf{P}_k & \mathbf{q}_k \\ \mathbf{P}_k^\top & \mathbf{S}_k & \mathbf{r}_k \\ \mathbf{q}_k^\top & \mathbf{r}_k^\top & z_k \end{bmatrix}. \quad (6)$$

In what follows, we investigate the dynamics of the training algorithm (5) at both the macroscopic and the microscopic levels. At the macroscopic level, by examining the cosines of the angles, we study how closely the model parameters $\mathbf{V}_k$, $\mathbf{w}_k$ associated with the generator and discriminator can align with the ground truth feature vectors, *i.e.*, the columns of $\mathbf{U}$. At the microscopic level, we study how the elements in the matrix $\mathbf{V}_k$ and the vector $\mathbf{w}_k$ evolve as a stochastic process. As our analysis will reveal, the mechanisms behind the two levels are different: the macroscopic dynamics is asymptotically deterministic whereas the microscopic dynamics stays stochastic even as $n \rightarrow \infty$.

3.1 Macroscopic dynamics

We first study the asymptotic dynamics of the macroscopic state $\mathbf{M}_k$. Our theoretical analysis is carried out under the following assumptions.

- (A.1) The sequences of $c_k \sim \mathcal{P}_c$ and $\tilde{c}_k \sim \mathcal{P}_{\tilde{c}}$ for $k = 0, 1, \dots$ are i.i.d. random variables with bounded moments of all orders, and $\{c_k\}$ is independent of $\{\tilde{c}_k\}$.
- (A.2) The sequences $\{\mathbf{a}_k\}$ and $\{\tilde{\mathbf{a}}_k\}$ for $k = 0, 1, \dots$ are both independent Gaussian vectors with zero mean and the covariance matrix $\mathbf{I}_n$. Moreover, $\{\mathbf{a}_k\}$, $\{\tilde{\mathbf{a}}_k\}$ are independent of $\{c_k\}$ and $\{\tilde{c}_k\}$.
- (A.3) The first-order derivative of $H(\cdot)$ and the derivatives up to fourth order of the functions $F(\hat{D}(\cdot))$ and $\tilde{F}(\hat{D}(\cdot))$ exist and they are also uniformly bounded.
- (A.4) Let $[\mathbf{U}, \mathbf{V}_0, \mathbf{w}_0]$ be the initial microscopic state. For $i = 1, 2, \dots, n$, we have $\mathbb{E}[\sum_{\ell=1}^d ([\mathbf{U}]_{i,\ell}^4 + [\mathbf{V}_0]_{i,\ell}^4 + [\mathbf{w}_0]_i^4)] \leq C/n^2$, where C is a constant not depending on n .
- (A.5) The initial macroscopic state $\mathbf{M}_0$ satisfies $\mathbb{E} \|\mathbf{M}_0 - \mathbf{M}_0^*\| \leq C/\sqrt{n}$, where $\mathbf{M}_0^*$ is a deterministic matrix and C is a constant not depending on n .

We provide a few remarks on the above assumptions. In Assumption (A.1), $\mathcal{P}_c$ and $\mathcal{P}_{\tilde{c}}$ can be different. For example, c is Gaussian, and $\tilde{c}$ is uniform on $[-1, 1]^d$. The assumption (A.2) can

be relaxed to non-Gaussian cases as long as all moments of $\mathbf{a}_k$ and $\tilde{\mathbf{a}}_k$ are bounded, but we use Gaussian assumption here to simplify the proof. The assumption (A.4) requires that the elements in the parameter matrix of real data $\mathbf{U}$ and initial microscopic state $\mathbf{X}_0$ are $\mathcal{O}(1/\sqrt{n})$ numbers. Intuitively, this assumption ensures that $\mathbf{U}$ and $\mathbf{X}_0$ are generic matrices with $\mathcal{O}(1)$ Frobenius norms (*i.e.*, not the matrices that most elements are zeros and only few elements are large numbers). The assumption (A.5) ensures that the initial macroscopic states converges to a deterministic value as the system size n goes to infinity. The following theorem proves that if the initial state is convergent, then the whole training process converges to a deterministic process as $n \rightarrow \infty$, which is characterized by an ODE.

Theorem 1. Fix $T > 0$. It holds under Assumptions (A.1)–(A.5) that

$$\max_{0 \leq k \leq nT} \mathbb{E} \|\mathbf{M}_k - \mathbf{M}(\frac{k}{n})\| \leq \frac{C(T)}{\sqrt{n}}, \quad (7)$$

where $C(T)$ is a constant that depends on T but not on n , and $\mathbf{M}(t) = \begin{bmatrix} \mathbf{I} & \mathbf{P}_t & \mathbf{q}_t \\ \mathbf{P}_t^\top & \mathbf{S}_t & \mathbf{r}_t \\ \mathbf{q}_t^\top & \mathbf{r}_t^\top & z_t \end{bmatrix} \in \mathbb{R}^{(2d+1) \times (2d+1)}$ is a deterministic function. Moreover, $\mathbf{M}(t)$ is the unique solution of the following ODE:

$$\begin{aligned} \frac{d}{dt} \mathbf{P}_t &= \tilde{\tau}(\mathbf{q}_t \tilde{\mathbf{g}}_t^\top + \mathbf{P}_t \mathbf{L}_t) \\ \frac{d}{dt} \mathbf{q}_t &= \tau(\mathbf{g}_t - \mathbf{P}_t \tilde{\mathbf{g}}_t + \mathbf{q}_t h_t) \\ \frac{d}{dt} \mathbf{r}_t &= \tau(\mathbf{P}_t^\top \mathbf{g}_t - \mathbf{S}_t \tilde{\mathbf{g}}_t + \mathbf{r}_t h_t) + \tilde{\tau}(z_t \tilde{\mathbf{g}}_t + \mathbf{L}_t \mathbf{r}_t) \\ \frac{d}{dt} \mathbf{S}_t &= \tilde{\tau}(\mathbf{r}_t \tilde{\mathbf{g}}_t^\top + \tilde{\mathbf{g}}_t \mathbf{r}_t^\top + \mathbf{S}_t \mathbf{L}_t + \mathbf{L}_t \mathbf{S}_t) \\ \frac{d}{dt} z_t &= 2\tau(\mathbf{q}_t^\top \mathbf{g}_t - \mathbf{r}_t^\top \tilde{\mathbf{g}}_t + z_t h_t) + \tau^2 b_t \end{aligned} \quad (8)$$

with the initial condition $\mathbf{M}(0) = \mathbf{M}_0^*$, where

$$\begin{aligned} \mathbf{g}_t &= \langle \mathbf{c} f(\mathbf{c}^\top \mathbf{q}_t + e\sqrt{z_t \eta_T}) \rangle_{\mathbf{c}, e}, \quad \tilde{\mathbf{g}}_t = \langle \tilde{\mathbf{c}} \tilde{f}(\tilde{\mathbf{c}}^\top \mathbf{r}_t + e\sqrt{z_t \eta_G}) \rangle_{\tilde{\mathbf{c}}, e}, \quad \mathbf{L}_t = -\lambda \text{diag}(H'(\mathbf{S}_t)) \\ h_t &= \langle f'(\mathbf{c}^\top \mathbf{q}_t + e\sqrt{z_t \eta_T}) \rangle_{\mathbf{c}, e} - \langle \tilde{f}'(\tilde{\mathbf{c}}^\top \mathbf{r}_t + e\sqrt{z_t \eta_G}) \rangle_{\tilde{\mathbf{c}}, e} - \lambda H'(z_t), \\ b_t &= \eta_T \langle f^2(\mathbf{c}^\top \mathbf{q}_t + e\sqrt{z_t \eta_T}) \rangle_{\mathbf{c}, e} + \eta_G \langle \tilde{f}^2(\tilde{\mathbf{c}}^\top \mathbf{r}_t + e\sqrt{z_t \eta_G}) \rangle_{\tilde{\mathbf{c}}, e}. \end{aligned} \quad (9)$$

The two functions $f, \tilde{f}$ stand for $f(x) = \frac{d}{dx} F(\hat{D}(x))$ and $\tilde{f}(x) = \frac{d}{dx} \tilde{F}(\hat{D}(x))$, and $f', \tilde{f}'$ and H' are derivatives of $f, \tilde{f}$ and H respectively. The two constants η_T and η_G are the strength of the noise in the true data model and the generator, respectively. The brackets $\langle \cdot \rangle_{\mathbf{c}, e}$ and $\langle \cdot \rangle_{\tilde{\mathbf{c}}, e}$ denote the averages over the random variables $\mathbf{c} \sim \mathcal{P}_{\mathbf{c}}$, $\tilde{\mathbf{c}} \sim \mathcal{P}_{\tilde{\mathbf{c}}}$, and $e \sim \mathcal{N}(0, 1)$, where $\mathcal{P}_{\mathbf{c}}$ and $\mathcal{P}_{\tilde{\mathbf{c}}}$ are the distributions involved in defining the generative model (1) and the generator (2).

This theorem implies that for each $k = \lfloor tn \rfloor$ for some $t \in [0, T]$, the macroscopic state $\mathbf{M}_k$ converges to a deterministic number $\mathbf{M}(t)$, and the convergence rate is $\mathcal{O}(1/\sqrt{n})$. The limiting ODE (8) for the macroscopic states involves $\mathcal{O}(d^2)$ variables, where d is the number of internal features often assumed to be a finite number that is much less than n . This ODE is essentially different from the ODE derived in the small-learning-rate limit [11–16], in which the number of variables is $\mathcal{O}(n)$.

The complete proof can be found in the Supplementary Materials. We briefly sketch the proof here. First, we note that $\mathbf{M}_k$ is a discrete-time stochastic process driven by the Markov chain $\mathbf{X}_k$. Then, we apply the martingale decomposition for $\mathbf{M}_k$ and get

$$\mathbf{M}_{k+1} - \mathbf{M}_k = \frac{1}{n} \phi(\mathbf{M}_k) + (\mathbf{M}_{k+1} - \mathbb{E}_k \mathbf{M}_{k+1}) + [\mathbb{E}_k \mathbf{M}_{k+1} - \mathbf{M}_k - \frac{1}{n} \phi(\mathbf{M}_k)],$$

where the matrix-valued function $\phi(\mathbf{M})$ represents the functions on the right hand sides of the ODE (8), and $\mathbb{E}_k$ denotes the conditional expectation given the state of the Markov chain $\mathbf{X}_k$. Finally, we show the martingale $\sum_{k'=0}^k (\mathbf{M}_{k'+1} - \mathbb{E}_{k'} \mathbf{M}_{k'})$ and the higher-order term $\mathbb{E}_k \mathbf{M}_{k+1} - \mathbf{M}_k - \frac{1}{n} \phi(\mathbf{M}_k)$ have no contribution when n goes to infinity.

Due to the limitation of our current proof, the constant $C(T)$ in (7) grows exponentially as T increases. This is not a problem for any finite T , but may cause some problem to study the long time behavior when $T \rightarrow \infty$. However, if we impose a sufficient large regularizer parameter λ to limit the norms of the microscopic weights $\mathbf{V}_k$ and $\mathbf{w}_k$, then the macroscopic state $\mathbf{M}_k$ is bounded

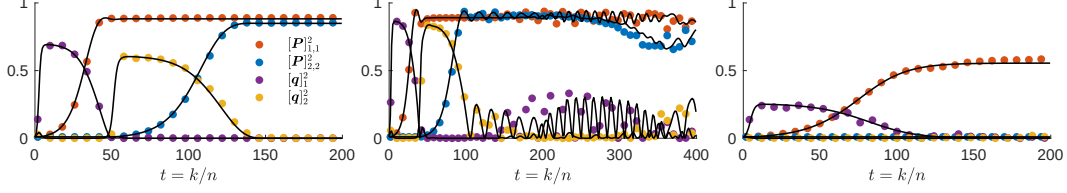


Figure 1: Macroscopic dynamics of the GAN with $d = 2$ features: $[P_k]_{i,j}$ is the cosine of the angle between i 'th column vector of the real feature matrix U_k and j 'th column vector of the generator's weight matrix V_k . Similarly, $[q_k]_i$ is the cosine of angle between i 'th column vector of U_k and the discriminator's weight vector w_k . Colored dots are results from experiments, and the curves tracing these dots are our theoretical prediction by the ODE (8). From the left to right, the variance of background noise is $\eta_T = \eta_G = 2, 1, 4$ respectively, and other parameters are the same. The left figure is an example of successful training, where two features (red and blue dots) are retrieved by the generator. The center figure shows an oscillating training. It happens when noise are weak. The right figure shows a mode collapsing state, in which only the first feature are estimated by the generator.

as $[M_k]_{i,j}^2 \leq [M_k]_{i,i}[M_k]_{j,j}$. In our experiments, $\lambda > 1$ is sufficient. In this case, the constant $C(T)$ is bounded not depending on T . In Example 1, when $\lambda \rightarrow \infty$, $[M_k]_{i,i} = 1$, and therefore $[M_k]_{i,j}^2 \leq 1$ and $C(T) \leq (2d+1)^2$, where the number of features d is considered a constant not growing with n . This justifies the fixed points analysis of the ODE as discussed in Section 4, which reflects the long-time training behavior. A better proof strategy to get rid of this dependence of T is also possible, *e.g.*, [26].

Numerical verification. We verify the theoretical prediction given by the ODE (8) via numerical simulations under the settings stated in Example 1. The results are shown in Figure 1. The number of features is $d = 2$, and c_k and $\tilde{c}_k$ are both Gaussian with zero mean and covariance $\text{diag}([5, 3])$. The dimension is $n = 5,000$, and the learning rates of the generator and discriminator are $\tilde{\tau} = 0.04$ and $\tau = 0.2$ respectively. After testing different noise strength $\eta_T = \eta_G = 2, 1, 4$, we have observed at least three nontrivial dynamical patterns: success, oscillating or mode collapsing. In all these experiments, our theoretical predictions match the actual trajectories of the macroscopic states pretty well.

Let us take a closer look at the successful case as shown in the left figure in Figure 1. The dynamics can be split into 4 stages. At the first stage, the discriminator learns the first feature of the true model. At this state, $[q_t]_1$ quickly increases. At the second stage, the generator starts to learn the first feature and the discriminator is deceived. At this stage, $[P_t]_{1,1}^2$ increases and $[q_t]_1^2$ decreases. Once the discriminator completely forgets the first feature as $[q_t]_1 \approx 0$, the third state begins. The discriminator starts to learn the second feature as $[q_t]_2^2$ increases. Then, at the last stage, the generator learns the second feature and the discriminator is fooled again. In this region, $[P_t]_{2,2}^2$ increases and $[q_t]_2^2$ decreases down to 0. Eventually, the generators learns both features and the discriminator is completely fooled. It ends up at a stationary state that $q_t = 0$ and P_t is nearly an identity matrix. Interestingly, this experiment shows that the generator learn features sequentially given a single-feature discriminator. This may be a reason why in practice, the discriminator's structure can be much simpler than the generator's.

3.2 Microscopic dynamics

In this section, we study how the elements in $X_k = [U, V_k, w_k]$ evolve during the training process. Instead of studying the trajectory of X_k , we study the evolution of the *empirical measure* of the microscopic states, which is defined as

$$\mu_k(\hat{u}, \hat{v}, \hat{w}) \stackrel{\text{def}}{=} \frac{1}{n} \sum_{i=1}^n \delta([\hat{u}^\top, \hat{v}^\top, \hat{w}] - \sqrt{n}[U]_{i,:}, [V_k]_{i,:}, [w_k]_i)$$

where $\delta(\cdot)$ is a Dirac measure on $\mathbb{R}^{2d+1}$ and $[U]_{i,:}, [V_k]_{i,:}$ are i th row of U and V_k respectively. The scaling factor $\sqrt{n}$ in the Dirac measures is introduced because $[U]_{i,\ell}, [V_k]_{i,\ell}$ and $[w_k]_i$ are $\mathcal{O}(1/\sqrt{n})$ quantities.

We next embed the discrete-time measure-valued stochastic process μ_k into a continuous-time process by defining $\mu_t^{(n)} \stackrel{\text{def}}{=} \mu_k(\hat{\mathbf{u}}, \hat{\mathbf{v}}, \hat{w})$ with $k = \lfloor nt \rfloor$. Following the general technical approach presented in [17], we can show that under the same assumptions as Theorem 1, given $T > 0$, the sequence of measure-valued process $\{\{\mu_t^{(n)}\}_{t \in [0, T]}\}_n$ converges weakly to a deterministic process $\{\mu_t\}_{t \in [0, T]}$. In addition, μ_t is the measure of the solution to the stochastic differential equation

$$\begin{aligned} d\hat{\mathbf{u}}_t &= 0 \\ d\hat{\mathbf{v}}_t &= \tilde{\tau}(\hat{w}_t \tilde{\mathbf{g}}_t + \mathbf{L}_t \hat{\mathbf{v}}_t) dt \\ d\hat{w}_t &= \tau(\hat{\mathbf{u}}_t^\top \mathbf{g}_t + \hat{\mathbf{v}}_t^\top \tilde{\mathbf{g}}_t + \hat{w}_t h_t) dt + \tau \sqrt{b_t} dB_t \end{aligned} \quad (10)$$

where $(\hat{\mathbf{u}}_0, \hat{\mathbf{v}}_0, \hat{w}_0) \sim \mu_0$; B_t is the standard Brownian motion. The functions $\mathbf{g}_t, \tilde{\mathbf{g}}_t, \mathbf{L}_t, h_t$ and b_t are defined in (9), in which the macroscopic quantities $\mathbf{P}_t, \mathbf{S}_t, \mathbf{q}_t, z_t, \mathbf{r}_t$ are computed as follows

$$\mathbf{P}_t = \langle \mu_t, \hat{\mathbf{u}} \hat{\mathbf{v}}^\top \rangle, \quad \mathbf{S}_t = \langle \mu_t, \hat{\mathbf{v}} \hat{\mathbf{v}}^\top \rangle, \quad \mathbf{q}_t = \langle \mu_t, \hat{\mathbf{u}} \hat{w} \rangle, \quad z_t = \langle \mu_t, \hat{w}^2 \rangle, \quad \mathbf{r}_t = \langle \mu_t, \hat{\mathbf{v}} \hat{w} \rangle, \quad (11)$$

where $\langle \mu_t, \cdot \rangle$ denotes the expectation with respect to the measure μ_t .

The SDE (10) shows the intuitive meaning of the functions defined in (9): $\mathbf{g}_t, \tilde{\mathbf{g}}_t, \mathbf{L}_t, h_t$ are drift coefficients of the SDE and b_t is the diffusion coefficient of the SDE. We also note that if one follows the analysis in the small-learning-rate limit [11–16], one will get an ODE for the microscopic states. Compared to our SDE formula, the diffusion term $\tau \sqrt{b_t} dB_t$ is missing in those works, and therefore the effect of the noise can not be analyzed.

Moreover, the deterministic measure μ_t is unique solution of the following PDE (given in its weak form): for any bounded smooth test function $\varphi(\hat{\mathbf{u}}, \hat{\mathbf{v}}, \hat{w})$,

$$\begin{aligned} \frac{d}{dt} \langle \mu_t, \varphi(\hat{\mathbf{u}}, \hat{\mathbf{v}}, \hat{w}) \rangle &= \\ \tilde{\tau} \langle \mu_t, (\hat{w} \tilde{\mathbf{g}}_t^\top + \hat{\mathbf{v}}^\top \mathbf{L}_t) \nabla_{\hat{\mathbf{v}}} \varphi \rangle &+ \tau \langle \mu_t, (\hat{\mathbf{u}}^\top \mathbf{g}_t - \hat{\mathbf{v}}^\top \tilde{\mathbf{g}}_t + h_t \hat{w}) \frac{\partial}{\partial \hat{w}} \varphi \rangle + \frac{\tau^2}{2} b_t \langle \mu_t, \frac{\partial^2}{\partial \hat{w}^2} \varphi \rangle \end{aligned} \quad (12)$$

where $\mathbf{q}_t, \mathbf{r}_t, \mathbf{S}_t$, and z_t are defined in (11), and the functions $\mathbf{g}_t, \tilde{\mathbf{g}}_t, b_t, h_t$ and $\mathbf{L}_t$ are defined in (9). We refer readers to [17] for a general framework for rigorously establishing the above scaling limit.

The connection between the microscopic and macroscopic dynamics can also be derived from the weak formulation of the PDE. Let φ being each element of $\hat{\mathbf{u}} \hat{\mathbf{v}}^\top, \hat{\mathbf{u}} \hat{w}, \hat{\mathbf{v}} \hat{w}, \hat{\mathbf{v}} \hat{\mathbf{v}}^\top, \hat{w}^2$, and substituting those φ into the PDE (12), we can derive the ODE (8). In the setting of this paper, the macroscopic dynamics enjoys a closed ODE: We can predict the macroscopic states without solving the PDE nor SDE at microscopic scale. However, in a more general setting, e.g. when we add a regularizer other than the L2 type, the ODE itself may not be closed. In that case, one has to solve the PDE directly.

Numerical verification. We verify the predictions given by the PDE (12) by setting $d = 1$ using a special choice of the $(n \times 1)$ -dimensional target feature matrix $\mathbf{U}$ whose elements are all $1/\sqrt{n}$ with $n = 10,000$. We also set the initial condition $\mu_0(\hat{\mathbf{v}}, \hat{w} | \hat{\mathbf{u}} = 1)$ to be a Gaussian distribution. (When $d = 1$, the macroscopic quantities $\mathbf{P}_t, \mathbf{q}_t, \mathbf{r}_t, \mathbf{S}_t$ reduce to scalars, so we remove their boldface here.) In this case, the PDE (12) admits a particularly simple analytical solution: at any time t , the solution $\mu_t(\hat{\mathbf{v}}, \hat{w} | \hat{\mathbf{u}} = 1)$ is a Gaussian distribution whose mean and covariance matrix are given by $\mathbb{E}_{\mu_t(\hat{\mathbf{v}}, \hat{w} | \hat{\mathbf{u}} = 1)} \begin{bmatrix} \hat{\mathbf{v}} \\ \hat{w} \end{bmatrix} = \begin{bmatrix} P_t \\ q_t \end{bmatrix}, \mathbb{E}_{\mu_t(\hat{\mathbf{v}}, \hat{w} | \hat{\mathbf{u}} = 1)} \begin{bmatrix} \hat{\mathbf{v}} \\ \hat{w} \end{bmatrix} \begin{bmatrix} \hat{\mathbf{v}} & \hat{w} \end{bmatrix} = \begin{bmatrix} S_t & r_t \\ r_t & z_t \end{bmatrix}$. Figure 2 overlays the contours of the probability distribution $\mu_t(\hat{\mathbf{v}}, \hat{w} | \hat{\mathbf{u}} = 1)$ at different times t over the point clouds of the actual experiment data $(\sqrt{n}[\mathbf{w}_k]_i, \sqrt{n}[\mathbf{V}_k]_{i,1})$. We can see that the theoretical prediction given by (12) has excellent agreement with simulation results.

4 Local Stability Analysis of the ODE for the Macroscopic States

In this section, we study how the parameters, such as the learning rates τ and $\tilde{\tau}$, noise strength η_G and η_T affect the training algorithm. We will focus on the concrete model as described in Example 1 so that we can have analytical solutions.

In order to further reduce the degrees of freedom of the ODE (8), we let the regularization parameter $\lambda \rightarrow \infty$. In this case, the vector $\mathbf{w}_k$ and all columns vectors of $\mathbf{V}_k$ are always normalized. Thus

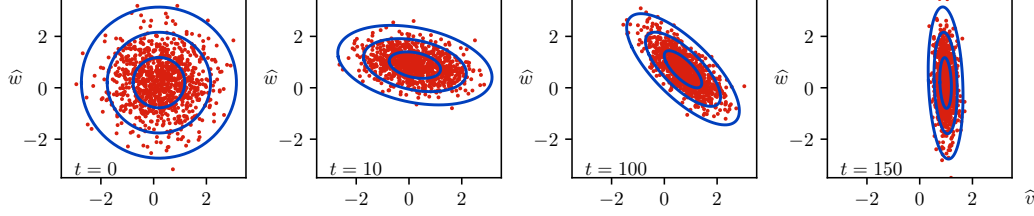


Figure 2: The evolution of the microscopic states at $t = 0, 10, 100$, and 150 . For each fixed t , the red points in the corresponding figure represent the values of $(\hat{v}, \hat{w}) = (\sqrt{n}[\mathbf{V}_k]_{i,1}, \sqrt{n}[\mathbf{w}_k]_i)$ for $i = 1, 2, \dots, n$, where $k = \lfloor nt \rfloor$. The blue ellipses illustrate the contours corresponding to one, two, and three standard deviations of the 2-D Gaussian distribution predicted by the PDE (12).

$z_k = 1$ and $[\mathbf{S}]_{i,i} = 1$. The macroscopic state is then described by $\mathbf{P}_k, \mathbf{q}_k, \mathbf{r}_k$ and off-diagonal terms of $\mathbf{S}_k$. Correspondingly, the ODE in Theorem 1 reduces to

$$\begin{cases} \frac{d}{dt} \mathbf{P}_t &= \tilde{\tau} (\mathbf{q}_t \mathbf{r}_t^\top \tilde{\Lambda} + \mathbf{P}_t \mathbf{L}_t) \\ \frac{d}{dt} \mathbf{q}_t &= \tau (\Lambda \mathbf{q}_t - \mathbf{P}_t \tilde{\Lambda} \mathbf{r}_t + h_t \mathbf{q}_t) \\ \frac{d}{dt} \mathbf{r}_t &= \tau (\mathbf{P}_t^\top \Lambda \mathbf{q}_t - \mathbf{S}_t \tilde{\Lambda} \mathbf{r}_t + h_t \mathbf{r}_t) + \tilde{\tau} (\tilde{\Lambda} + \mathbf{L}_t) \mathbf{r}_t \\ \frac{d}{dt} \mathbf{S}_t &= \tilde{\tau} (\mathbf{r}_t \mathbf{r}_t^\top \tilde{\Lambda} + \tilde{\Lambda} \mathbf{r}_t \mathbf{r}_t^\top + \mathbf{S}_t \mathbf{L}_t + \mathbf{L}_t \mathbf{S}_t) \end{cases} \quad (13)$$

where Λ and $\tilde{\Lambda}$ are the covariance matrices of the distributions P_c and $P_{\tilde{c}}$, respectively; and

$$h_t = (1 - \frac{\tau \eta_G}{2}) \mathbf{r}_t^\top \tilde{\Lambda} \mathbf{r}_t - (1 + \frac{\tau \eta_T}{2}) \mathbf{q}_t^\top \Lambda \mathbf{q}_t - \tau \frac{\eta_G + \eta_T^2}{2}, \quad \mathbf{L}_t = -\text{diag}(\mathbf{r}_t \mathbf{r}_t^\top \tilde{\Lambda}), \quad (14)$$

in which η_T and η_G are the variance of noise in the true data model and generator, respectively. The derivation from the ODE (8) to (13) is presented in the Supplementary Materials.

Next, we discuss under what conditions, the GAN can reach a desirable training state by studying local stability of a particular type of fixed points of the ODE (13). The perfect estimation of the generator corresponds to $\mathbf{P}_t$ being an identity matrix (up to a permutation of rows and columns). A complete fail state relates to $\mathbf{P} = \mathbf{0}$. Furthermore, It is easy to verify that if $\mathbf{q}_t = \mathbf{r}_t = \mathbf{0}$, the ODE (13) will be stable for any $\mathbf{P}_t = \mathbf{P}$.

Claim 1. *The macroscopic states $\mathbf{P}_t, \mathbf{q} = \mathbf{r} = \mathbf{0}$ for all valid $\mathbf{P}_t$ are always the fixed points of the ODE (13). Furthermore, a sufficient condition that the perfect estimation state $\mathbf{P}_t = \mathbf{I}, \mathbf{q} = \mathbf{r} = \mathbf{0}$ is locally stable and the failed state $\mathbf{P}_t = \mathbf{0}, \mathbf{q} = \mathbf{r} = \mathbf{0}$ is unstable if*

$$\frac{1}{2} \max_{\ell} \{ \Lambda_{\ell} - \tilde{\Lambda}_{\ell} + \alpha \tilde{\Lambda}_{\ell} \} \leq \tau \bar{\eta}^2 < \min_{\ell} \Lambda_{\ell}, \quad (15)$$

where $\alpha = \frac{\tilde{\tau}}{\tau}$, $\bar{\eta}^2 = \frac{1}{2}(\eta_T^2 + \eta_G^2)$, and $\Lambda_{\ell} = [\Lambda]_{\ell,\ell}$, $\tilde{\Lambda}_{\ell} = [\tilde{\Lambda}]_{\ell,\ell}$.

The proof can be found in the Supplementary Materials. If the right inequality in (15) is violated, any feature ℓ with the signal-to-noise ratio $[\Lambda]_{\ell,\ell} < \tau \bar{\eta}^2$ is not learned by the generator resulting *mode collapsing*. The right figure in Figure 1 demonstrates this situations, where only one of the two features is recovered. If the left inequality in (15) is violated, the training processes can be trapped in an *oscillation phase*. This phenomenon is shown in the middle figure in Figure 1. This result indicates that proper background noise can help to avoid oscillation and stabilize the training process. In fact, the trick of injecting additional noise has been used in practice to train multi-layer GANs [27]. To our best knowledge, our paper is the first theoretical study on why noise can have such a positive effect via a dynamic perspective.

In experiments, the training is not ended at the perfect recovery point due to the presence of the noise but converges at another fixed point nearby. This is because the perfect state is marginally stable, as the Jacobian matrix always has zero eigenvalues. It indicates that there are other locally stable fixed points near $\mathbf{P} = \mathbf{I}$. In fact, all points in the hyper-rectangle region satisfying $\mathbf{q} = \mathbf{r} = \mathbf{0}$ and $|p_{\ell}^*| \leq |[\mathbf{P}]_{\ell,\ell}| \leq 1, \forall \ell = 1, 2, \dots, d$ are locally stable for some critical p_{ℓ}^* . In the matched case when $\Lambda_{\ell} = \tilde{\Lambda}_{\ell}$, we have $p_{\ell}^* = [(\Lambda_{\ell} - \tau \bar{\eta}^2)(\tilde{\Lambda}_{\ell} + \tau \bar{\eta}^2 - \alpha \tilde{\Lambda}_{\ell}) / (\Lambda_{\ell} \tilde{\Lambda}_{\ell})]^{1/2}$, $\alpha = \frac{\tilde{\tau}}{\tau}$ and $\bar{\eta}^2 =$

$\frac{1}{2}(\eta_T^2 + \eta_G^2)$. Starting from a point near the origin, numerical solution of the ODE shows the training processes are ended up at the corner of this hyper-rectangle, *i.e.*, $\mathbf{P}^* = \text{diag}(\{p_\ell^*, \ell = 1, 2, \dots, d\})$. In the small-learning rate limit $\tau \rightarrow 0$ and the learning rate ratio $\alpha \rightarrow 0$, we get the perfect recovery $\mathbf{P}^* = \mathbf{I}$. The limit $\tau \rightarrow 0, \alpha \rightarrow 0$ was studied in the small-learning-rate analysis with the two-time scaling [15], and the result is consistent, but our analysis includes the situations with finite τ and α .

In addition, we provide a phase diagram analysis in a single-feature case $d = 1$ in the Supplementary Materials. All possible fixed points in this case are enumerated and their local stability is analyzed. This helps us understand the successful recovery condition (15), which is the intersection of the informative phases that each feature can be recovered individually.

5 Conclusion

We present a simple high-dimensional model for GAN with an exactly analyzable training process. Using the tool of scaling limits of stochastic processes, we show that the macroscopic state associated with the training process converges to a deterministic process characterized as the unique solution of an ODE, whereas the microscopic state remains stochastic described by an SDE, whose time-varying probability measure is described by a limiting PDE.

Indeed, it is a common picture in statistical physics that the macroscopic states of large systems tend to converge to deterministic values due to self-averaging. These notions, especially the mean-field dynamics, have been applied to analyzing neural networks both in shallow [19, 20] and deep models [28]. However, this mean-field regime was not considered in previous analyses of GAN. For example, a series of recent works *e.g.*, [11–16] considers a different scaling regime where the learning rate goes to zero but the system dimension n stays fixed. In that regime, the microscopic dynamics are deterministic even with the presence of the microscopic noise. In contrast, we study the regime where the learning rate is fixed but the dimension $n \rightarrow \infty$. This setting allows us to quantify the effect of training noise in the learning dynamics.

In this paper, we only consider a linear generator with a latent variable $\tilde{\mathbf{c}}$ drawn from a fixed distribution $\mathcal{P}_{\tilde{\mathbf{c}}}$, but our analysis can be extended to a more complex non-linear model with a learnable latent-variable distribution. Specifically, in order to compute derivatives w.r.t. $\mathcal{P}_{\tilde{\mathbf{c}}}$, the latent variable $\tilde{\mathbf{c}} \sim \mathcal{P}_{\tilde{\mathbf{c}}}$ should be reparameterized by a deterministic function $\tilde{\mathbf{c}} = f(\mathbf{z}; \boldsymbol{\theta})$, where $\boldsymbol{\theta}$ is a learnable parameter and $\mathbf{z}$ is a random variable drawn from a simple and fixed distribution. For example, a Gaussian mixture with L equal-probability modes can be parameterized by $\tilde{\mathbf{c}} = \sum_{\ell=1}^L (\boldsymbol{\mu}_\ell + \boldsymbol{\Sigma}_\ell \boldsymbol{\epsilon}_\ell) \beta_\ell$, where $\boldsymbol{\mu}_\ell$ and $\boldsymbol{\Sigma}_\ell$ are two learnable parameters representing the mean and covariance of the ℓ th mode respectively, and $\boldsymbol{\epsilon} \sim \mathcal{N}(0, \mathbf{I})$; β_ℓ is a random indicator variable where only one β_ℓ for $\ell = 1, 2, \dots, L$ is 1 and the others are 0. In practice, $f(\mathbf{z}; \boldsymbol{\theta})$ is implemented by a multilayer neural network. Our analysis can be naturally extended to analyzing this model as long as the dimensions of $\tilde{\mathbf{c}}$ and $\boldsymbol{\theta}$ keep finite when the data dimension n goes to infinity. More challenging situations, where the dimension of $\boldsymbol{\theta}$ is proportional to n , will be explored in future works.

Although our analysis is carried out in the asymptotic setting, numerical experiments show that our theoretical predictions can accurately capture the actual performance of the training algorithm at moderate dimensions. Our analysis also reveals several different phases of the training process that highly depend on the choice of the learning rates and noise strength. The analysis reveals a condition on the learning rates and the strength of noise to have successful training. Violating this condition results either oscillation or mode collapsing. Despite its simplicity, the proposed model of GAN provides a new perspective and some insights for the study of more realistic models and more involved training algorithms.

Acknowledgments This work was supported by the US Army Research Office under contract W911NF-16-1-0265 and by the US National Science Foundation under grants CCF-1319140, CCF-1718698, and CCF-1910410.

References

- [1] I. Goodfellow, J. Pouget-Abadie, M. Mirza, B. Xu, D. Warde-Farley, S. Ozair, A. Courville, and Y. Bengio, “Generative Adversarial Nets,” in *Advances in Neural Information Processing System*, 2014, pp. 2672–2680.
- [2] M. Arjovsky, S. Chintala, and L. Bottou, “Wasserstein Generative Adversarial Networks,” *Proceedings of The 34th International Conference on Machine Learning*, pp. 1–32, 2017.
- [3] M. Lucic, K. Kurach, M. Michalski, S. Gelly, and O. Bousquet, “Are gans created equal? a large-scale study,” in *Advances in neural information processing systems*, 2018, pp. 698–707.
- [4] C. Ledig, L. Theis, F. Huszar, J. Caballero, A. Cunningham, A. Acosta, A. Aitken, A. Tejani, J. Totz, Z. Wang, and W. Shi, “Photo-Realistic Single Image Super-Resolution Using a Generative Adversarial Network,” in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2017, pp. 4681–4690.
- [5] P. Isola, J.-Y. Zhu, T. Zhou, and A. A. Efros, “Image-to-Image Translation with Conditional Adversarial Networks,” in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2017, pp. 1125–1134.
- [6] S. Reed, Z. Akata, X. Yan, L. Logeswaran, B. Schiele, and H. Lee, “Generative Adversarial Text to Image Synthesis,” *33rd International Conference on Machine Learning*, pp. 1060–1069, 2016.
- [7] S. Arora, R. Ge, Y. Liang, and Y. Zhang, “Generalization and Equilibrium in Generative Adversarial Nets,” in *International Conference on Machine Learning*, 2017, pp. 224–232.
- [8] M. Arjovsky and L. Bottou, “Towards Principled Methods for Training Generative Adversarial Networks,” *arXiv preprint arXiv:1701.04862*, 2017.
- [9] S. Feizi, C. Suh, F. Xia, and D. Tse, “Understanding GANs: the LQG Setting,” *arXiv preprint arXiv:1710.10793*, 2017.
- [10] J. Li, A. Madry, J. Peebles, and L. Schmidt, “Towards Understanding the Dynamics of Generative Adversarial Networks,” *arXiv preprint arXiv:1706.09884*, 2017.
- [11] L. Mescheder, A. Geiger, and S. Nowozin, “Which training methods for gans do actually converge?” in *International Conference on Machine Learning*, 2018, pp. 3478–3487.
- [12] L. Mescheder, S. Nowozin, and A. Geiger, “The numerics of GANs,” in *Advances in Neural Information Processing Systems*, 2017, pp. 1823–1833.
- [13] V. Nagarajan and J. Z. Kolter, “Gradient descent GAN optimization is locally stable,” in *Advances in Neural Information and Processing Systems*, 2017, pp. 5591–5600.
- [14] K. Roth, A. Lucchi, S. Nowozin, and T. Hofmann, “Stabilizing Training of Generative Adversarial Networks through Regularization,” in *Advances in Neural Information Processing Systems*, 2017, pp. 2015–2025.
- [15] M. Heusel, H. Ramsauer, T. Unterthiner, B. Nessler, and S. Hochreiter, “GANs Trained by a Two Time-Scale Update Rule Converge to a Local Nash Equilibrium,” in *Advances in Neural Information Processing Systems*, 2017, pp. 6629–6640.
- [16] E. V. Mazumdar, M. I. Jordan, and S. S. Sastry, “On finding local nash equilibria (and only local nash equilibria) in zero-sum games,” *arXiv preprint arXiv:1901.00838*, 2019.
- [17] C. Wang, J. Mattingly, and Y. M. Lu, “Scaling Limit: Exact and Tractable Analysis of Online Learning Algorithms with Applications to Regularized Regression and PCA,” *arXiv preprint arXiv:1712.04332*, 2017.
- [18] G. O. Roberts, A. Gelman, and W. R. Gilks, “Weak convergence and optimal scaling of random walk Metropolis algorithms,” *Annals of Applied Probability*, vol. 7, no. 1, pp. 110–120, 1997.
- [19] D. Saad and S. A. Solla, “Exact Solution for On-Line Learning in Multilayer Neural Networks,” *Phys. Rev. Lett.*, vol. 74, no. 21, pp. 4337–4340, 1995.
- [20] M. Biehl and H. Schwarze, “Learning by on-line gradient descent,” *Journal of Physics A*, vol. 28, no. 3, pp. 643–656, 1995.
- [21] C. Wang and Y. M. Lu, “Online Learning for Sparse PCA in High Dimensions: Exact Dynamics and Phase Transitions,” in *Information Theory Workshop (ITW), 2016 IEEE*, 2016, pp. 186–190.
- [22] C. Wang, Y. C. Eldar, and Y. M. Lu, “Subspace estimation from incomplete observations: A high-dimensional analysis,” *IEEE Journal of Selected Topics in Signal Processing*, vol. 12, no. 6, pp. 1240–1252, Dec 2018.
- [23] C. Wang and Y. M. Lu, “The Scaling Limit of High-Dimensional Online Independent Component Analysis,” in *Advances in Neural Information Processing Systems*, 2017, pp. 6641–6650.
- [24] S. Mei, A. Montanari, and P.-M. Nguyen, “A Mean Field View of the Landscape of Two-Layers Neural Networks,” *arXiv preprint*, p. arXiv:1804.06561, 2018.
- [25] I. Johnstone and A. Lu, “On consistency and sparsity for principal components analysis in high dimensions,” *Journal of the American Statistical Association*, vol. 104, no. 486, pp. 682–693, 2009.
- [26] B. Jourdain, T. Lelièvre, and B. Miasojedow, “Optimal scaling for the transient phase of Metropolis Hastings algorithms: The longtime behavior,” *Bernoulli*, vol. 20, no. 4, pp. 1930–1978, 2014.
- [27] C. K. Sønderby, J. Caballero, L. Theis, W. Shi, and F. Huszár, “Amortised map inference for image super-resolution,” *International Conference on Learning Representations*, 2017.
- [28] P.-M. Nguyen, “Mean field limit of the learning dynamics of multilayer neural networks,” *arXiv preprint arXiv:1902.02880*, 2019.