

An Interpolation Perspective on Neural Network Generalization

Jean-Michel Morel

Joint work with Jin GUO and Roy HE

Lingnan University

FAU, Erlangen

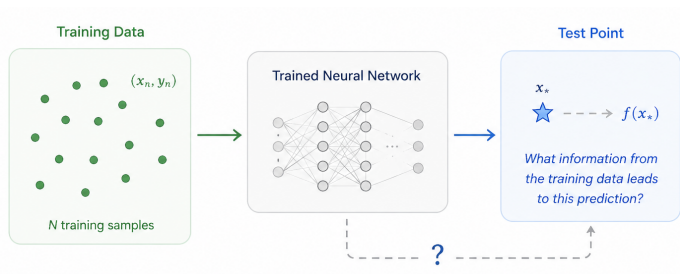
June 21, 2026

Content

- 1 Motivation
- 2 Domingos Formula
- 3 Stochastic Domingos Theorem
- 4 Generalization
- 5 Path Kernel Perspective for Generative Models
- 6 Second-order Domingos Theorem
- 7 Summary

Motivation

How does a trained neural network generalize?



We consider any parametric machine trained by gradient descent on a dataset of examples of (input, output) pairs. Gradient descent is applied to a training loss expressing how well a network fits the observed samples. The question of generalization is how a trained network behaves at an unseen input, namely how it relates the input to the training inputs.

The Domingos formula gives an explicit answer to this question.

Domingos Formula: A Path-Kernel View

Consider the general training setup:

- Dataset $\{(x_n, y_n^*)\}_{n=1}^N$, neural network $f(x, \Theta) : \mathbb{R}^p \times \mathbb{R}^d \rightarrow \mathbb{R}^m$.
- Loss: $L(\Theta) = \frac{1}{N} \sum_{n=1}^N \ell(f(x_n, \Theta), y_n^*)$.
- Training dynamics: $\Theta_k = \Theta_{k-1} - \eta \nabla_{\Theta} L(\Theta_{k-1})$.
- The associated continuous gradient flow, or learning dynamics is defined by

$$\dot{\Theta}_t = -\nabla L(\Theta_t)$$

We define the *path kernel* associated to the learning dynamics by

$$K_t(x, x') := \langle \nabla_{\Theta} f(x, \Theta_t), \nabla_{\Theta} f(x', \Theta_t) \rangle.$$

Domingos Formula: A Path-Kernel View

Consider the general training setup:

- Dataset $\{(x_n, y_n^*)\}_{n=1}^N$, neural network $f(x, \Theta) : \mathbb{R}^p \times \mathbb{R}^d \rightarrow \mathbb{R}^m$.
- Loss: $L(\Theta) = \frac{1}{N} \sum_{n=1}^N \ell(f(x_n, \Theta), y_n^*)$.
- Training dynamics: $\Theta_k = \Theta_{k-1} - \eta \nabla_{\Theta} L(\Theta_{k-1})$.

When the learning rate $\eta \rightarrow 0$ and $\eta k \rightarrow T$, the Domingos formula¹ is

$$\lim_{\eta \rightarrow 0} f(x, \Theta_T) = f(x, \Theta_0) - \frac{1}{N} \sum_{n=1}^N \int_0^T \underbrace{\langle \nabla_{\Theta} f(x, \Theta_t), \nabla_{\Theta} f(x_n, \Theta_t) \rangle}_{\substack{\text{interaction weight;} \\ \text{path kernel}}} \underbrace{\frac{\partial \ell}{\partial y_n}(x_n, \Theta_t)}_{\text{training correction}} dt.$$

The path kernel measures how strongly the test point interacts with the training sample in the tangent feature space at time t .

¹Pedro Domingos. Every model learned by gradient descent is approximately a kernel machine. arXiv preprint arXiv:2012.00152, 2020.

Domingos Formula: an interpolation formula

For a quadratic loss, the Domingos formula becomes

$$f(x, \Theta_T) = f(x, \Theta_0) - \frac{1}{N} \sum_{n=1}^N \int_0^T \underbrace{\langle \nabla_{\Theta} f(x, \Theta_t), \nabla_{\Theta} f(x_n, \Theta_t) \rangle}_{\text{interaction weight // path kernel}} \underbrace{(f(x_n, \Theta_t) - y_n^*)}_{\text{training residual}} dt.$$

This formula makes it apparent that the output of the learning machine for the input x is a weighted average along the path of the residuals of all training samples, weighted by the path kernel. The path kernel appears like a weight that measures the interaction between the input x and the other samples x_n

Domingos Theorem

Theorem (Informal) Let $f(\cdot, \Theta)$ be a differentiable model trained by gradient descent $\dot{\Theta}_t = -\nabla L(\Theta_t)$ on the data $\{(x_n, y_n^*)\}_{n=1}^N$, using the empirical loss $L(\Theta) = \frac{1}{N} \sum_{n=1}^N \ell(f(x_n, \Theta), y_n^*)$. Then for any input $x \in \mathbb{R}^p$, the model output at time T satisfies

$$f(x, \Theta_T) = f(x, \Theta_0) - \frac{1}{N} \sum_{n=1}^N \int_0^T K_t(x, x_n) \frac{\partial \ell}{\partial f}(f(x_n, \Theta_t), y_n^*) dt, \quad (1)$$

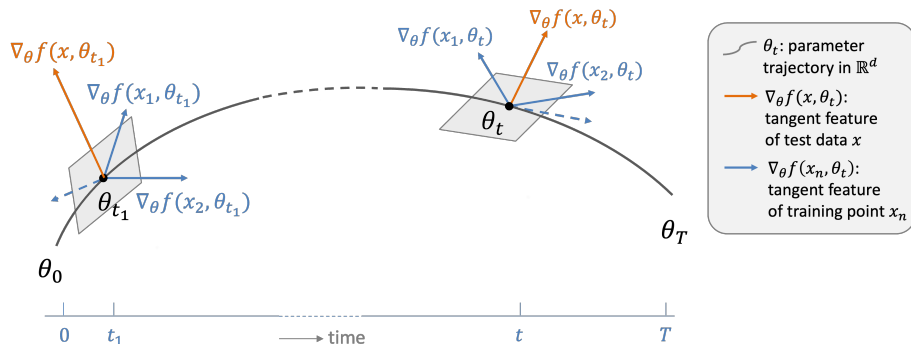
Proof. Let $\Theta_t = (\Theta_t^1, \dots, \Theta_t^d) \in \mathbb{R}^d$ denote the parameter of the neural network $f(\cdot, \Theta)$. By the chain rule and the gradient flow $\dot{\Theta}_t = -\nabla L(\Theta_t)$, we have

$$\begin{aligned} \frac{df}{dt}(x, \Theta_t) &= - \sum_{i=1}^d \frac{\partial f}{\partial \Theta_t^i}(x, \Theta_t) \frac{\partial L(\Theta_t)}{\partial \Theta_t^i} \\ &= - \sum_{i=1}^d \frac{\partial f}{\partial \Theta_t^i}(x, \Theta_t) \left(\frac{1}{N} \sum_{n=1}^N \frac{\partial f}{\partial \Theta_t^i}(x_n, \Theta_t) \frac{\partial \ell}{\partial f}(f(x_n, \Theta_t), y_n^*) \right) \\ &= - \frac{1}{N} \sum_{n=1}^N \sum_{i=1}^d \frac{\partial f}{\partial \Theta_t^i}(x, \Theta_t) \frac{\partial f}{\partial \Theta_t^i}(x_n, \Theta_t) \frac{\partial \ell}{\partial f}(f(x_n, \Theta_t), y_n^*) \\ &= - \frac{1}{N} \sum_{n=1}^N \langle \nabla_{\Theta} f(x, \Theta_t), \nabla_{\Theta} f(x_n, \Theta_t) \rangle \frac{\partial \ell}{\partial f}(f(x_n, \Theta_t), y_n^*). \end{aligned}$$

Integrating both sides from 0 to T gives (1).

Domingos Theorem

$$f(x, \Theta_T) = f(x, \Theta_0) - \frac{1}{N} \sum_{n=1}^N \int_0^T \underbrace{\langle \nabla_{\Theta} f(x, \Theta_t), \nabla_{\Theta} f(x_n, \Theta_t) \rangle}_{\text{interaction weight}} \underbrace{(f(x_n, \Theta_t) - y_n^*)}_{\text{training residual}} dt.$$



Is Domingos formula general?

- What if training is made by stochastic gradient descent: Real training is done by stochastic optimizers using random mini batches.
- Is there a path kernel representation for stochastic optimizers?
- How does the Domingos formula depend on the kind of stochastic optimization (SGDM, RMSprop, Adam)?
- What does this imply for generalization?

Continuous Time Approximation of the discrete SGD

The simplest stochastic gradient descent (SGD) is defined by

$$\Theta_{k+1} = \Theta_k - \eta \nabla_{\Theta} L_{\mathcal{B}_k}(\Theta_k),$$

where $\mathcal{B}_k$ is a random uniform mini-batch with fixed cardinality and

$$L_{\mathcal{B}_k}(\Theta_k) = \frac{1}{|\mathcal{B}_k|} \sum_{n \in \mathcal{B}_k} \ell(f(x_n, \Theta_k), y_n^*).$$

To generalize Domingos formula, we need a continuous model and an asymptotic formulation of the learning dynamics as a stochastic differential equation.

(Several propositions for a continuous SDE [Mandt, 2015; Jastrzebski, 2018, Zhu, 2019])

Stochastic Domingos Theorem

Lemma (First-order weak approximation of SGD²) Let $T > 0$, $\eta \in (0, 1 \wedge T)$, and set $K = \lfloor T/\eta \rfloor$. Let $\{\Theta_k\}_{k \geq 0}$ be the SGD iterates. Define $\{Z_t : t \in [0, T]\}$ as the solution of the SDE

$$dZ_t = -\nabla L(Z_t) dt + \sqrt{\eta} \Sigma_{|\mathcal{B}|}^{1/2}(Z_t) dW_t, \quad Z_0 = \Theta_0, \quad (2)$$

where

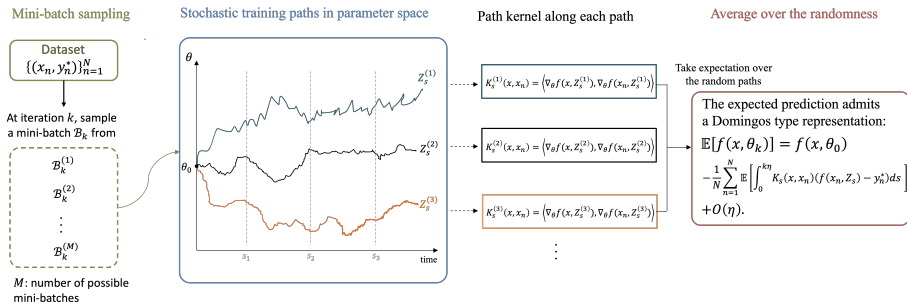
$$\Sigma_{|\mathcal{B}|}(Z) = \mathbb{E}_{\mathcal{B}}[(\nabla L_{\mathcal{B}}(Z) - \nabla L(Z))(\nabla L_{\mathcal{B}}(Z) - \nabla L(Z))^{\top}].$$

Then $\{Z_t\}$ is a first-order weak approximation of the SGD iterates; for any smooth test function g ,

$$\max_{k=0, \dots, K} |\mathbb{E}[g(\Theta_k)] - \mathbb{E}[g(Z_{k\eta})]| = O(\eta). \quad (3)$$

²Qianxiao Li, Cheng Tai, and E Weinan. Stochastic modified equations and dynamics of stochastic gradient algorithms I: Mathematical foundations. *Journal of Machine Learning Research*, 20(40):1–47, 2019.

Stochastic Domingos Theorem



Theorem (Informal) Consider a learning model $y = f(x, \Theta)$. Assume that Θ is learned from a dataset $\{(x_n, y_n^*)\}_{n=1}^N$ by SGD with learning rate η . Then

$$\mathbb{E}[f(x, \Theta_k)] = f(x, \Theta_0) - \frac{1}{N} \sum_{n=1}^N \mathbb{E} \left[\int_0^{k\eta} \underbrace{K_s(x, x_n)}_{\text{interaction weight}} \underbrace{\frac{\partial \ell}{\partial f}(f(x_n, Z_s), y_n^*)}_{\text{training correction}} ds \right] + O(\eta), \quad (4)$$

where $K_s(x, x') := \langle \nabla_{\Theta} f(x, Z_s), \nabla_{\Theta} f(x', Z_s) \rangle$ is the stochastic path kernel.

Stochastic Domingos' Theorem (SGDM)

Stochastic Gradient Descent with Momentum

Theorem (Informal) Consider a learning model $y = f(x, \Theta)$. Assume that Θ is learned from the dataset $\{(x_n, y_n^*)\}_{n=1}^N$ by SGDM with learning rate η and momentum weight $\beta \in [0, 1)$. Let $\mu = (1 - \beta)/\eta$. Then

$$\mathbb{E}[f(x, \Theta_k)] = f(x, \Theta_0) - \frac{1}{N} \sum_{n=1}^N \mathbb{E} \left[\int_0^{k\eta} \int_0^t \underbrace{K_{t,s}(x, x_n)}_{\text{interaction with memory decay}} \underbrace{\frac{\partial \ell}{\partial f}(f(x_n, Z_s), y_n^*)}_{\text{training correction}} ds dt \right] + O(\eta), \quad (5)$$

where the path kernel has become

$$K_{t,s}(x, x') := \langle \nabla_{\Theta} f(x, Z_t), e^{-(t-s)\mu} \nabla_{\Theta} f(x', Z_s) \rangle.$$

When μ is large (weak momentum), the memory kernel collapses and the result turns to the SGD case.

Stochastic Domingos' Theorem (RMSprop)

Theorem (Informal) Consider a learning model $y = f(x, \Theta)$. Assume that Θ is learned from the dataset $\{(x_n, y_n^*)\}_{n=1}^N$ by RMSprop (Root Mean Square Propagation) with learning rate η , second-moment decay rate $\beta \in [0, 1)$, and smoothing constant $\varepsilon > 0$. Let $\mu = (1 - \beta)/\eta$ and suppose that $\mu = O(1)$ as $\eta \rightarrow 0$. Then

$$\mathbb{E}[f(x, \Theta_k)] = f(x, \Theta_0) - \frac{1}{N} \sum_{n=1}^N \mathbb{E} \left[\int_0^{k\eta} \underbrace{K_t^{P_t}(x, x_n)}_{\text{preconditioned interaction weight}} \underbrace{\frac{\partial \ell}{\partial f}(f(x_n, Z_t), y_n^*)}_{\text{training correction}} dt \right] + O(\eta), \quad (6)$$

where path kernel $K_t^{P_t}(x, x') = \langle \nabla_{\Theta} f(x, Z_t), \nabla_{\Theta} f(x', Z_t) \rangle_{P_t^{-1}}$ and

$$P_t = \text{diag} \left(\mu \int_0^t e^{-\mu(t-s)} [(\nabla L(Z_s))^2 + \text{diag}(\Sigma_{|\mathcal{B}|}(Z_s))] ds \right) + \varepsilon \mathbf{I}_d, \\ \Sigma_{|\mathcal{B}|}(Z) = \mathbb{E}_{\mathcal{B}} [(\nabla L_{\mathcal{B}}(Z) - \nabla L(Z))(\nabla L_{\mathcal{B}}(Z) - \nabla L(Z))^{\top}].$$

Stochastic Domingos' Theorem (Adam)

Consider a learning model $y = f(x, \Theta)$. Assume that Θ is learned from the dataset $\{(x_n, y_n^*)\}_{n=1}^N$ by Adam with learning rate η , first-moment coefficient $\beta_1 \in [0, 1)$, second-moment coefficient $\beta_2 \in [0, 1)$, and smoothing constant $\varepsilon > 0$. Let $c_1 = (1 - \beta_1)/\eta$, $c_2 = (1 - \beta_2)/\eta$, and suppose that $c_1 = O(\eta^{-\xi})$ for some $\xi \in (0, 1)$ and $c_2 = O(1)$ as $\eta \rightarrow 0$. Then

$$\begin{aligned} \mathbb{E}[f(x, \Theta_k)] &= f(x, \Theta_0) - c_1 \beta_1 \frac{1}{N} \sum_{n=1}^N \mathbb{E} \left[\int_0^{k\eta} \int_0^t K_{t,s}^{P_t}(x, x_n) \frac{\partial \ell}{\partial f}(f(x_n, Z_s), y_n^*) ds dt \right] \\ &\quad - c_1(1 - \beta_1) \frac{1}{N} \sum_{n=1}^N \mathbb{E} \left[\int_0^{k\eta} K_{t,t}^{P_t}(x, x_n) \frac{\partial \ell}{\partial f}(f(x_n, Z_t), y_n^*) dt \right] + O(\sqrt{\eta}), \end{aligned} \quad (7)$$

where stochastic path kernel

$$K_{t,s}^{P_t}(x, x') = \langle \nabla_{\Theta} f(x, \Theta_t), \frac{\sqrt{1 - e^{-c_2 t}}}{1 - e^{-c_1 t}} e^{-c_1(t-s)} \nabla_{\Theta} f(x', \Theta_t) \rangle_{P_t^{-1}}$$

and $P_t = c_2^{1/2} \text{diag} \left(\int_0^t e^{-c_2(t-s)} [\text{diag}(\Sigma_{|B|}(Z_s)) + (\nabla L(Z_s))^2] ds \right)^{1/2} + \varepsilon \sqrt{1 - e^{-c_2 t}} \mathbf{I}_d$.

Evolution of gradient kernel during training

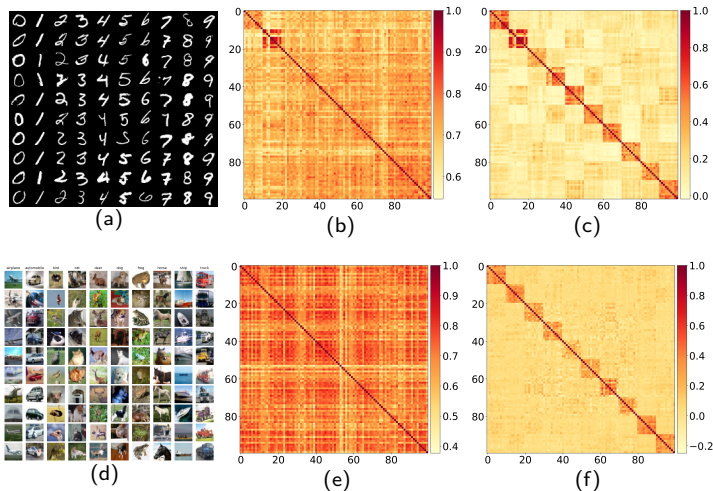


Figure: Normalized path kernel $\hat{K}_t(x_i, x_j) := K_t(x_i, x_j) / \sqrt{K_t(x_i, x_i) K_t(x_j, x_j)}$ on MNIST and CIFAR-10. (a,d) Example samples. (b,e) Kernels at initialization. (c,f) Kernels after training.

Extrapolation failure in the gradient kernel view

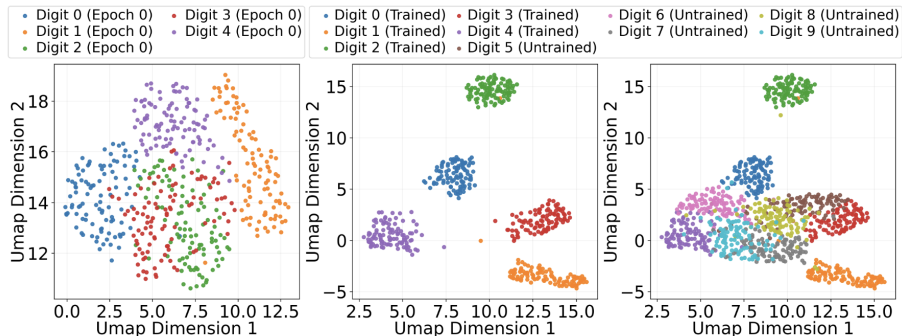


Figure: UMAP visualization of MNIST digits in the tangent feature space $\{\nabla_{\Theta} f(x_n, \Theta_k)\}_{n=1}^N$. **Left:** Initialization (0-th epoch) on digits 0–4. **Middle:** After training on digits 0–4 (200-th epoch). **Right:** Tangent-feature visualization of all digits: while trained digits (0–4) remain well separated, unseen digits (5–9) fail to form distinct clusters.

Extrapolation failure in the gradient kernel view

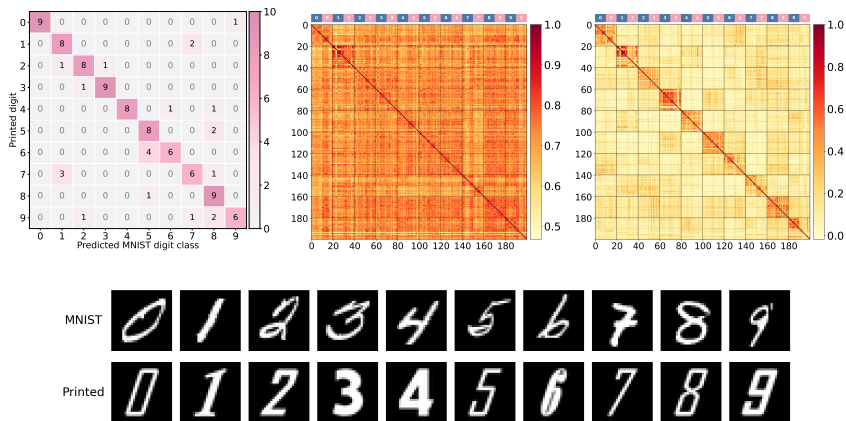


Figure: Normalized gradient kernel for out-of-domain printed digits. Top: normalized gradient kernels at initialization (a) and after training (b), and the printed-digit confusion matrix (c). Bottom: MNIST and printed digit examples. Blue and pink bars indicate MNIST and printed blocks, respectively.

Extrapolation failure in the gradient kernel view

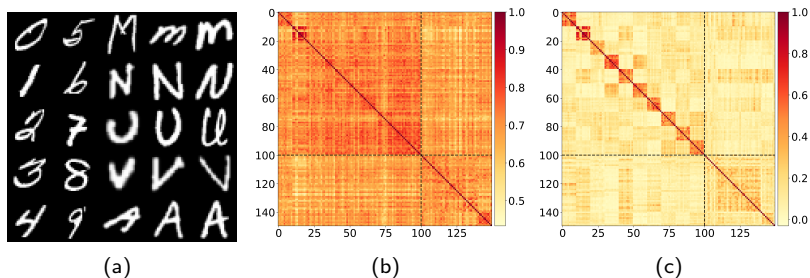


Figure: Normalized gradient kernel for out-of-domain inputs (unseen letters). Left: MNIST digits and unseen handwritten letters *M*, *N*, *U*, *V*, *A*. Middle and right: normalized gradient kernels at initialization and after training.

RKHS reminder

A positive definite kernel $K : \mathcal{X} \times \mathcal{X} \rightarrow \mathbb{R}$ generates a Reproducing kernel Hilbert space (RKHS) $\mathcal{H}_K$ with inner product

$$\langle K(x, \cdot), K(y, \cdot) \rangle_{\mathcal{H}_K} = K(x, y).$$

For every $f \in \mathcal{H}_K$, the reproducing property gives:

$$f(x) = \langle f, K(x, \cdot) \rangle_{\mathcal{H}_K}, \forall x \in \mathcal{X}.$$

Given a data distribution μ , $x' \sim \mu$, define the integral kernel operator

$$[T_K g](x) = \int_{\mathcal{X}} K(x, x') g(x') d\mu(x').$$

If $g(x')$ represents the training residual (correction) at point x' , then $K(x, x')g(x')$ measures how much the correction at x' contributes to the prediction at x .

By integrating over all x' , we obtain the total correction received at x .

Generalization Error and Kernel Null Space

Let $\mathcal{K}_t$ be the integral operator induced by stochastic path kernel $K_t(x, x') = \langle \nabla_{\Theta} f(x, \Theta_t), \nabla_{\Theta} f(x', \Theta_t) \rangle$, defined by

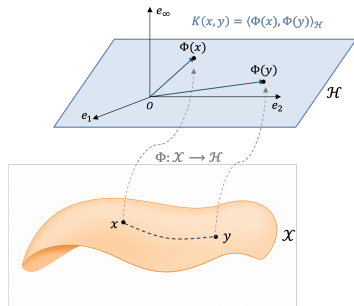
$$[\mathcal{K}_t g](x) = \int_{\mathbb{R}^p} K_t(x, x') g(x') d\mathcal{P}(x'),$$

where $\mathcal{P}$ is the marginal distribution of x on $\mathbb{R}^p$. The stochastic path kernel induces a trajectory-dependent RKHS

$$\mathcal{H}_t := \mathcal{H}_{\mathcal{K}_t}.$$

Mercer's theorem gives

$K_t(x, x') = \sum_{i=1}^{\infty} \lambda_i \phi_i(x) \phi_i(x')$ and $\ker(\mathcal{K}_t) = \text{span}\{\phi_i : \lambda_i = 0\}$ collects the unlearnable directions.



Generalization Error and Kernel Null Space

Theorem (Informal) Under mild assumptions, let f^* be the target function and suppose the loss is the mean-squared error. Define the terminal residual $\Delta_T(x) := f^*(x) - f(x, Z_T)$ and the generalisation error

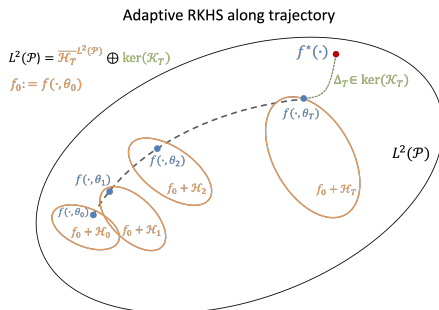
$$\mathcal{E}_T := \|\Delta_T\|_{L^2(\mathbb{R}^p, \mathcal{P})}^2 = \int_{\mathbb{R}^p} (f^*(x) - f(x, Z_T))^2 d\mathcal{P}(x).$$

Then, as $\eta \rightarrow 0$,

$$\forall \varepsilon > 0, \quad \mathbb{P}(\Delta_T \in \mathcal{N}_\varepsilon(\ker(\mathcal{K}_T))) \rightarrow 1,$$

where $\mathcal{N}_\varepsilon(\cdot)$ denotes the ε -neighbourhood in $L^2(\mathbb{R}^p, \mathcal{P})$. In particular, the generalisation error $\mathcal{E}_T$ is governed by the component of the residual lying in $\ker(\mathcal{K}_T)$.

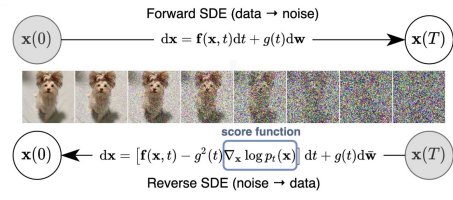
Generalization Error and Kernel Null Space



At each time t , the stochastic path kernel K_t induces an RKHS $\mathcal{H}_t$.

Path Kernel Perspective for DDPM

(Denoising Diffusion Probabilistic Model)



Corollary (Informal) Let $\{(\mathbf{x}_{t,n}, \gamma(t)), \varepsilon_t\}_{t=1, n=1}^{T, N}$ denote the training set, where T is the total number of timesteps and N is the number of training samples. Let $\{Z_s\}$ denote the weak first-order approximation of the SGD iterates. Then, for any query $(\mathbf{x}, \gamma(\tau))$, the expected noise prediction satisfies

$$\begin{aligned} \mathbb{E}[\varepsilon_{\Theta}((\mathbf{x}, \gamma(\tau)), \Theta_k)] &= \varepsilon_{\Theta}((\mathbf{x}, \gamma(\tau)), \Theta_0) \\ &- \frac{2}{N \cdot T} \sum_{t=1}^T \sum_{n=1}^N \mathbb{E} \left[\int_0^{k\eta} \|\varepsilon_{\Theta}((\mathbf{x}_{t,n}, \gamma(t)), Z_s) - \varepsilon_t\| K_s((\mathbf{x}, \gamma(\tau)), (\mathbf{x}_{t,n}, \gamma(t))) ds \right] \\ &+ O(\eta), \end{aligned}$$

where K_s is the stochastic gradient kernel of the noise predictor ε_{Θ} .

Path Kernel Perspective for Generative Models

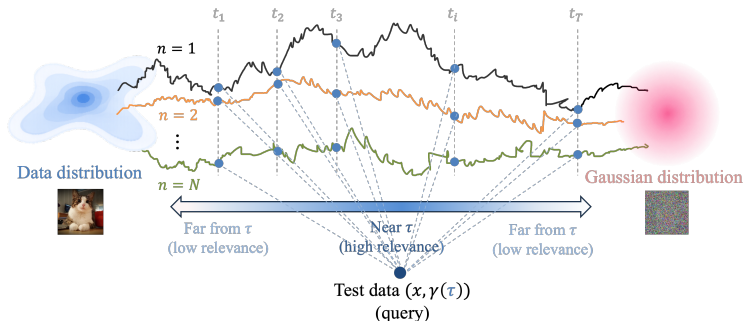
Path memory over all samples and noise levels

stored denoising residual

$$r(t, n, Z_s) := \left\| \varepsilon_{\theta} \left((x_{t,n}, \gamma(t)), Z_s \right) - \varepsilon_t \right\|$$

$t \in \{1, 2, \dots, T\}$: noise levels (time partition)

$n \in \{1, 2, \dots, N\}$: trajectories (samples)



Expected generated output

$$\mathbb{E}[\varepsilon_{\theta}((x, \gamma(\tau)), \theta_k)] = \varepsilon_{\theta}((x, \gamma(\tau)), \theta_0) - \frac{1}{NT} \sum_{t=1}^T \sum_{n=1}^N \mathbb{E} \left[\int_0^{kn} \underbrace{r(t, n, Z_s)}_{\text{noise-prediction residual}} \underbrace{K_s((x, \gamma(\tau)), (x_{t,n}, \gamma(t)))}_{\text{interaction weights}} ds \right] + O(\eta).$$

Path kernel perspective for GAN

Generative adversarial networks (GANs) learn a generator $G(\cdot, \cdot) : \mathcal{Z} \times \mathbb{R}^P \rightarrow \mathcal{X}$ and a discriminator $D(\cdot) : \mathcal{X} \rightarrow (0, 1)$ via the minimax objective

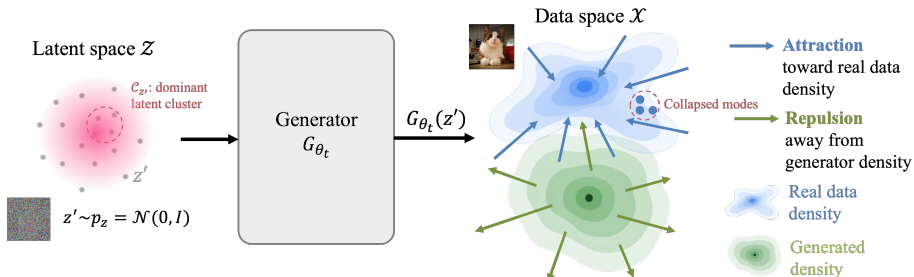
$$\min_{\Theta} \max_D \left\{ \mathbb{E}_{x \sim p_{\text{data}}} \log D(x) + \mathbb{E}_{z \sim p_Z} \log (1 - D(G(z, \Theta))) \right\}.$$

Corollary (Informal) Consider a generator $G_{\Theta} := G(\cdot, \Theta)$ trained by SGD with learning rate η on latent samples $z' \sim p_Z$. Assume that the discriminator is fully optimized at each iteration and treating D^* as fixed when differentiating with respect to Θ . Let $\{\bar{\Theta}_t\}$ denote the weak first-order approximation of the SGD iterates. Then, when $\nabla_y \log p_{\text{data}}$ and $\nabla_y \log p_{G_t}$ are well-defined, for any query z , the expected generator output satisfies

$$\begin{aligned} \mathbb{E}[G_{\Theta_k}(z)] &= G_{\Theta_0}(z) - \mathbb{E} \left[\int_0^{k\eta} \mathbb{E}_{z' \sim p(z)} \left[K_t(z, z') \frac{\partial}{\partial G} \ell(D^*(G_{\bar{\Theta}_t}(z'))) \right] dt \right] + O(\eta) \\ &= G_{\Theta_0}(z) + \mathbb{E} \left[\int_0^{k\eta} \mathbb{E}_{z' \sim p(z)} \left[K_t(z, z') D^*(G_{\bar{\Theta}_t}(z')) (\nabla_y \log p_{\text{data}}(y) - \nabla_y \log p_{G_{\bar{\Theta}_t}}(y)) \Big|_{y=G_{\bar{\Theta}_t}(z')} \right] dt \right] \\ &\quad + O(\eta), \end{aligned} \tag{8}$$

where K_t is the stochastic gradient kernel of the generator $G_{\bar{\Theta}_t}$, and $\ell(D(G_{\bar{\Theta}_t}(z))) = \log(1 - D(G_{\bar{\Theta}_t}(z)))$.

Path Kernel Perspective for GAN



discriminator-induced correction field

$$v_t(z', \theta_t) := D^* \left(G_{\theta_t}(z') \right) \left(\nabla_y \log p_{data}(y) - \nabla_y \log p_{G_{\theta_t}}(y) \right) \Big|_{y = G_{\theta_t}(z')}$$

Expected generated output

$$\mathbb{E}[G(z, \theta_k)] = G(z, \theta_0) + \mathbb{E} \int_0^{k\eta} \mathbb{E}_{z' \sim p(z)} \left[\underbrace{K_t(z, z')}_{\text{interaction weights}} \underbrace{v_t(z', \theta_t)}_{\text{correction field}} \right] dt + O(\eta)$$

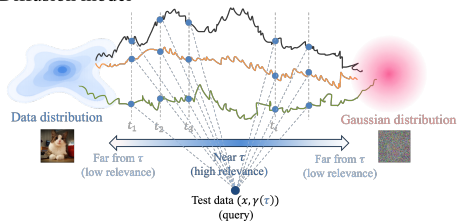
Mode collapse amplification:

A subset of latent samples $\mathcal{C}_{z'}$ dominate the correction field $v_t(z', \theta_t)$.

$\Rightarrow$ Mode collapse is reinforced during training.

Path Kernel Correction for Generative Model

Diffusion model



DM's correction(residual):

$$r(t, n, Z_s) := \left\| \varepsilon_{\theta} \left((x_{t,n}, y(t)), Z_s \right) - \varepsilon_t \right\|$$

GAN's correction

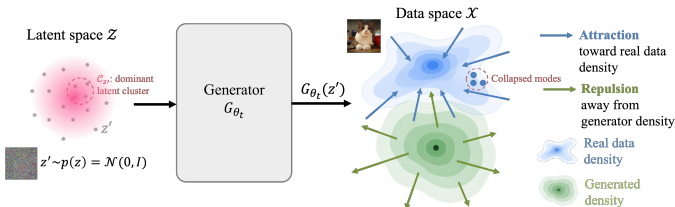
$$v_t(z', \theta_t) := D^* \left(G_{\theta_t}(z') \right) \left(\nabla_y \log p_{data}(y) - \nabla_y \log p_{G_{\theta_t}}(y) \right) \Big|_{y = G_{\theta_t}(z')}$$

Difference:

DM retrieves supervised residuals from a full noise-indexed memory.

GAN propagates distribution-mismatch corrections from the current generated support.

GAN



What Is Missing from First-Order Path Kernels?

- Domingos' formula gives a first-order explanation of prediction:
 $\text{prediction} = \text{initialization} + \text{accumulated tangent feature interpolation}.$
- However, practical training is not purely first-order.
 - ▶ The finite learning rate, the curvature of the loss landscape, and the randomness of mini batch sampling introduce higher-order corrections to the prediction.

Can we make these corrections explicit in the interpolation formula?
Does a higher-order description still preserve the interpolation structure?

Second-Order Interpolation: Gradient Descent

Theorem (Informal) Consider a differentiable model $f(\cdot, \Theta)$ trained on a set $\{(x_n, y_n^*)\}_{n=1}^N$ via gradient descent with learning rate $\eta > 0$. The model output satisfies

$$\begin{aligned} f(x, \Theta_K) = & f(x, \Theta_0) - \frac{1}{N} \sum_{n=1}^N \int_0^T \underbrace{K_t(x, x_n)}_{\text{path kernel}} \underbrace{\frac{\partial \ell}{\partial f}(f(x_n, \Theta_t), y_n^*)}_{\text{training correction}} dt \\ & - \frac{\eta}{2N} \sum_{n=1}^N \int_0^T \underbrace{K_t^{\text{cur}}(x, x_n)}_{\text{curvature-induced kernel}} \underbrace{\frac{\partial \ell}{\partial f}(f(x_n, \Theta_t), y_n^*)}_{\text{training correction}} dt + O(\eta^2), \end{aligned}$$

where path kernel $K_t(x, x') := \langle \nabla_{\Theta} f(x, \Theta_t), \nabla_{\Theta} f(x', \Theta_t) \rangle$, curvature-induced kernel $K_t^{\text{cur}}(x, x') = \langle \nabla_{\Theta} f(x, \Theta_t), \nabla_{\Theta} f(x', \Theta_t) \rangle_{\nabla_{\Theta}^2 L(\Theta_t)}$.

Remark. Under a learning-rate schedule $\eta_k = \eta_{\min} + \frac{1}{2}(\eta_{\max} - \eta_{\min})\left(1 + \cos\left(\frac{k}{K}\pi\right)\right)$, $k = 0, \dots, K$, the prediction satisfies

$$\begin{aligned} f(x, \Theta_K) = & f(x, \Theta_0) - \frac{1}{N} \sum_{n=1}^N \int_0^T \left(w(t) - \frac{\eta}{2} \dot{w}(t) \right) K_t(x, x_n) \frac{\partial \ell}{\partial f}(f(x_n, \Theta_t), y_n^*) dt \\ & - \frac{\eta}{2N} \sum_{n=1}^N \int_0^T w(t)^2 K_t^{\text{cur}}(x, x_n) \frac{\partial \ell}{\partial f}(f(x_n, \Theta_t), y_n^*) dt + O(\eta^2), \end{aligned}$$

with the scalar weight $w(t) = 1 + \frac{1}{2}(\kappa - 1)(1 + \cos(\frac{\pi t}{T}))$, $\kappa = \eta_{\max}/\eta_{\min}$, $t \in [0, T]$.

Second-Order Interpolation: Gradient Descent

Proof. Let $\Theta_t \in \mathbb{R}^d$ denote the parameter trajectory of the model $f(\cdot, \Theta)$. Following the standard modified equation for gradient descent, we approximate the discrete update $\Theta_{k+1} = \Theta_k - \eta \nabla L(\Theta_k)$ by continuous-time dynamics

$$\dot{\Theta}_t = -\nabla_{\Theta} L(\Theta_t) - \frac{\eta}{2} \nabla_{\Theta}^2 L(\Theta_t) \nabla_{\Theta} L(\Theta_t) + O(\eta^2). \quad (9)$$

By the chain rule and (9), we obtain

$$\begin{aligned} \frac{df}{dt}(x, \Theta_t) &= \nabla_{\Theta} f(x, \Theta_t)^{\top} \frac{d\Theta_t}{dt} \\ &= \nabla_{\Theta} f(x, \Theta_t)^{\top} \left(-\nabla_{\Theta} L(\Theta_t) - \frac{\eta}{2} \nabla_{\Theta}^2 L(\Theta_t) \nabla_{\Theta} L(\Theta_t) + O(\eta^2) \right) \\ &= -\frac{1}{N} \sum_{n=1}^N \langle \nabla_{\Theta} f(x, \Theta_t), \nabla_{\Theta} f(x_n, \Theta_t) \rangle \frac{\partial \ell}{\partial f}(f(x_n, \Theta_t), y_n^*) \\ &\quad - \frac{\eta}{2N} \sum_{n=1}^N \langle \nabla_{\Theta} f(x, \Theta_t), \nabla_{\Theta} f(x_n, \Theta_t) \rangle_{\nabla_{\Theta}^2 L(\Theta_t)} \frac{\partial \ell}{\partial f}(f(x_n, \Theta_t), y_n^*) + O(\eta^2), \end{aligned}$$

where $\langle \nabla_{\Theta} f(x, \Theta_t), \nabla_{\Theta} f(x_n, \Theta_t) \rangle_{\nabla_{\Theta}^2 L(\Theta_t)} = \nabla_{\Theta} f(x, \Theta_t)^{\top} \nabla_{\Theta}^2 L(\Theta_t) \nabla_{\Theta} f(x_n, \Theta_t)$.

Integrating both sides from 0 to T yields the result.

Second-Order Interpolation: SGD

Compared with GD, SGD introduces an additional source of randomness: the mini-batch gradient noise.

Theorem (Informal) Consider a learning model $y = f(x, \Theta)$ trained by SGD with learning rate η . In the sense of a second-order weak approximation of SGD, we have

$$\begin{aligned}\mathbb{E}[f(x, \Theta_K)] &= f(x, \Theta_0) - \frac{1}{N} \sum_{n=1}^N \mathbb{E} \left[\int_0^T K_s(x, x_n) \frac{\partial \ell}{\partial f}(f(x_n, Z_s), y_n^*) ds \right] \\ &\quad - \frac{\eta}{2N} \sum_{n=1}^N \mathbb{E} \left[\int_0^T K_s^{\text{cur}}(x, x_n) \frac{\partial \ell}{\partial f}(f(x_n, Z_s), y_n^*) ds \right] \\ &\quad + \frac{\eta}{2} \mathbb{E} \left[\int_0^T \underbrace{\text{Tr}(\nabla_{\Theta}^2 f(x, Z_s) \Sigma_{|\mathcal{B}|}(Z_s))}_{\text{sampling-induced term}} ds \right] + O(\eta^2),\end{aligned}$$

where $K_t(x, x') := \langle \nabla_{\Theta} f(x, \Theta_t), \nabla_{\Theta} f(x', \Theta_t) \rangle$, curvature-induced kernel $K_t^{\text{cur}}(x, x') = \langle \nabla_{\Theta} f(x, \Theta_t), \nabla_{\Theta} f(x', \Theta_t) \rangle_{\nabla_{\Theta}^2 L(\Theta_t)}$.

This sampling-induced term is specific to SGD and depends on the mini-batch sampling noise, it reflects the interaction between mini-batch randomness and Hessian of the model output.

Second-Order Interpolation: SGDM

Theorem (Informal) Consider a learning model $y = f(x, \Theta)$ learned from a dataset $\{(x_n, y_n^*)\}_{n=1}^N$ using stochastic gradient descent with momentum (SGDM). In the sense of a second-order weak approximation of SGDM, we have

$$\begin{aligned} \mathbb{E}[f(x, \Theta_K)] &= f(x, \Theta_0) - \frac{1}{N} \sum_{n=1}^N \mathbb{E} \left[\int_0^T \int_0^t K_{t,s}(x, x_n) \frac{\partial \ell}{\partial f}(f(x_n, Z_s), y_n^*) ds dt \right] \\ &\quad - \frac{\eta}{2N} \sum_{n=1}^N \mathbb{E} \left[\int_0^T \int_0^t K_{t,s}^{\text{cur}}(x, x_n) \frac{\partial \ell}{\partial f}(f(x_n, Z_s), y_n^*) ds dt \right] \\ &\quad - \frac{\eta}{2N} \sum_{n=1}^N \mathbb{E} \left[\int_0^T K_t(x, x_n) \frac{\partial \ell}{\partial f}(f(x_n, Z_t), y_n^*) dt \right] \\ &\quad + \underbrace{\sqrt{\eta} \mathbb{E} \left[\int_0^T \int_0^t e^{-\mu(t-s)} \text{Tr} \left((D_s Z_t)^\top \nabla_{\Theta}^2 f(x, Z_t) \Sigma_{|\mathcal{B}|}^{1/2}(Z_s) \right) ds dt \right]}_{\text{sampling-induced term}} + O(\eta^2), \end{aligned}$$

where $\mu = \frac{1-\beta}{\eta}$ is fixed, β is momentum weight, and

$$\begin{aligned} K_{t,s}(x, x_n) &:= \left\langle \nabla_{\Theta} f(x, Z_t), e^{-\mu(t-s)} \left(1 - \frac{\mu^2}{2}(t-s)\eta \right) \nabla_{\Theta} f(x_n, Z_s) \right\rangle, \\ K_{t,s}^{\text{cur}}(x, x_n) &:= \left\langle \nabla_{\Theta} f(x, Z_t), e^{-\mu(t-s)} \nabla_{\Theta} f(x_n, Z_s) \right\rangle \int_s^t \nabla_{\Theta}^2 L(Z_r) dr. \end{aligned}$$

Here, D_s denotes the Malliavin derivative with respect to the driving Brownian motion at time s .

Fluctuation around the interpolation formula

Higher-order terms still preserve the interpolation structure.

They refine it by adding curvature-weighted and sampling-induced interpolation corrections.

The formula gives the mean prediction, but how far is the realized prediction from it?

A simplified form of [Proposition 4.1](#)³ gives, for any $\delta \in (0, 1)$,

$$\mathbb{P}\left(\|f(x, \Theta_K) - \mathbb{E}[f(x, \Theta_K)]\| \leq S_{|\mathcal{B}|} G_f(x) \sqrt{\eta \frac{e^{2\tilde{K}_\eta T} - 1}{\tilde{K}_\eta} \log \frac{2}{\delta}}\right) \geq 1 - \delta.$$

Here $S_{|\mathcal{B}|}$ is an upper bound for the mini-batch noise covariance, $G_f(x) > 0$ bounds $\|\nabla_{\Theta} f(x, z)\|$ on $\mathcal{U}$, and $\tilde{K}_\eta$ collects the regularity constants of the drift and diffusion terms.

³Jin Guo, Roy Y. He, and Jean-Michel Morel. Second-Order Path Kernel Interpolation Formulas in Machine Learning. arXiv preprint arXiv:2606.07495, 2026.

Scaling order verification

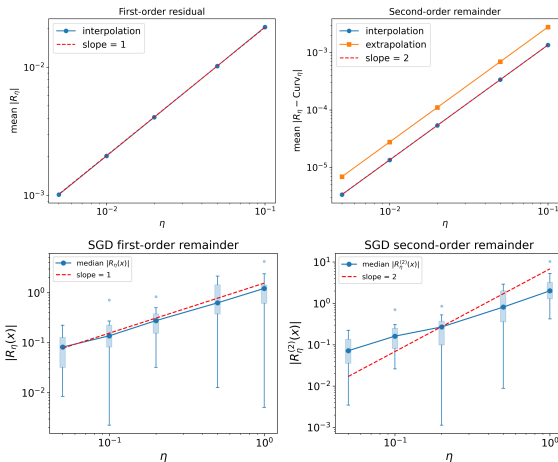


Figure: Scaling verification of the output expansion. Top: GD remainders show first-order $O(\eta)$ and second-order $O(\eta^2)$ scaling after subtracting the corresponding path-kernel and curvature terms. Bottom: SGD remainders show the same empirical scaling across 20 independent runs.

- **Training as memory.** Neural networks store residuals along the parameter path. Predictions are obtained by retrieving these residuals through path-kernel interaction weights. Thus, the network prediction is a trajectory-dependent interpolation of the training data.
- **Stochasticity changes the memory weights.** Mini-batch noise does not destroy the path-kernel representation. Different optimizers change the weights of training corrections.
- **Generalization through adaptive RKHS.** The stochastic path kernel induces an adaptive RKHS, and the generalization error lies in the null space of the associated kernel operator.
- **Generative models through the same lens.** Diffusion models and GANs both admit a path-kernel correction form.

THANK YOU!

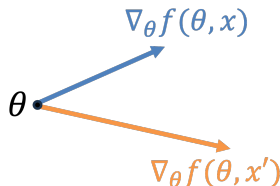
Neural Tangent Kernel

NTK gives a static kernel view of neural network prediction.

- Neural tangent kernel(NTK)^[1]:

$$K_{\theta}^{\text{NTK}}(x, x') = \langle \nabla_{\theta} f(\theta, x), \nabla_{\theta} f(\theta, x') \rangle$$

measures the similarity between data points x, x' by comparing their gradients.



- Under infinite width limit and NTK initialization, the NTK view is static.

$$K_{\theta_0}^{\text{NTK}}(x, x') \approx K_{\theta_{\infty}}^{\text{NTK}}(x, x').$$

The kernel is (approximately) fixed at initialization.

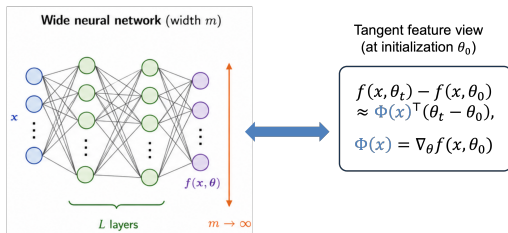
[1] Jacot, A., Gabriel, F. and Hongler, C., 2018. Neural tangent kernel: Convergence and generalization in neural networks. Advances in neural information processing systems, 31.

Neural Tangent Kernel

- Wide neural networks are linear in the parameter space

$$f(x, \theta_t) = f(x, \theta_0) + \langle \nabla_{\theta} f(x, \theta_0), \theta_t - \theta_0 \rangle + O\left(\frac{1}{m}\right), m: \text{width of NN.}$$

The NTK regime views neural network training through a nearly fixed linearized model around initialization.



Useful for analyzing infinite-width NNs.

But finite width training can move beyond this linearized regime.