

This talk is about

On Lazy training in Diff. PDE.
by Chizat - Ouyallon - Bach

three big things today

1) Any parametric model
trained by a gradient flow
has an induced function
evolution governed by
a time-dependent kernel.
This is exact.

2) If the tangent kernel

Stays nearly constant,
training is a kernel
gradient descent in
a fixed RKHS.

3) Lazy training is a
Where the NTK phenomenon
exists, yet ⁱⁿ many practical settings,
Good performance correlates
with leaving the NTK regime.

Setup: Let $f_\theta: \mathbb{R}^d \rightarrow \mathbb{R}$ (diff.)
 $\theta \in \mathbb{R}^p$

Given $(x_i, y_i)_{i=1}^n$, define

$$L(\theta) = \frac{1}{2} \sum_{i=1}^n (f_\theta(x_i) - y_i)^2$$

G-flow is

$$\dot{\theta} = -\nabla_\theta L(\theta(t))$$

Notation:

$$f_t(x) := f_{\theta(t)}(x)$$

$$\vec{f}(t) := (f_t(x_1) \dots f_t(x_n))^T \in \mathbb{R}^n$$

$$\vec{y} := (y_1, \dots, y_n)^T$$

Calculus time

$$\bar{\theta} = -\nabla_{\theta} \mathcal{L} = -\sum_{i=1}^n (f_{\theta}(x_i) - y_i) \nabla_{\theta} f_{\theta}(x_i)$$

$$\partial_t f_t = \nabla_{\theta} f_{\theta(t)}(x) \cdot \dot{\theta}(t)$$

$$= -\sum_{i=1}^n \langle \nabla_{\theta} f_t(x), \nabla_{\theta} f(x_i) \rangle (f_t(x_i) - y_i)$$

$$= -\sum_{i=1}^n \underbrace{\Theta_{\theta(t)}(x, x_i)}_{\text{Universal Kernel}} (f_t(x_i) - y_i)$$

Universal Kernel

Define

$$\Theta(t) \in \mathbb{R}^{n \times n}$$

$$\Theta(t)_{ij} := \Theta_{\text{act}}(x_{ij})$$

Then

$$\dot{\vec{f}} = -\Theta(t)(\vec{f}(t) - \vec{y})$$

$$\text{Let } r(t) = \vec{f}(t) - \vec{y}$$

Then

$$\dot{r}(t) = -\Theta(t)r(t)$$

Q: How does Θ behave
when we use neural nets?

Q2: What would be
the easiest Θ to
work with? (A: constant
+ P.D.)

Assume $\Theta \equiv \Theta_0$ constant, P.D.

$$\dot{r}(t) = \Theta_0 r(t)$$

$$r(0) = \vec{f}(0) - \vec{y}$$

$$r(t) = e^{-t\Theta_0} r(0)$$

$$\Rightarrow \vec{f}(t) = \vec{y} + e^{-t\Theta_0} (\vec{f}(0) - \vec{y})$$

$\rightarrow \vec{y}$ as $t \rightarrow \infty$.

Perfect. Now time
for a kernel trick

Define

$$K_0(x) := \begin{pmatrix} \Theta_0(x, x_1), \dots, \\ \Theta_0(x, x_n) \end{pmatrix}^T$$

(Still assuming $\Theta(0) = \Theta(t) \dots$)

test point ODE (just rewriting universal kernel eq)

$$\begin{aligned} \partial_t f_t(x) &= -K_0(x)^T r(t) \\ &= -K_0(x)^T e^{-t\Theta_0} r(0) \end{aligned}$$

$$\begin{aligned} \rightarrow f_t(x) &= f_0(x) \\ &\quad - K_0(x)^T \int_0^t e^{-s\Theta_0} ds r(0) \end{aligned}$$

$$= f_0(x) - K_0(x)^T \Theta_0^{-1} (I - e^{-t\Theta_0}) r(0)$$

$\xrightarrow{t \rightarrow \infty}$

$$= f_0(x) - K_0(x)^T \Theta_0^{-1} (\hat{f}(0) - \bar{y})$$

Imagine you had $f_0(x) \equiv 0$

$$f_\infty(x) = K_0(x)^T \Theta_0^{-1} y$$

the representer form of
kernel interpolation on
the training set.

End of (2)

This is a complete
theory! IF only $\Theta(t)$ constant.

Start (3) :

Let's look at linearization
the typical way.

$$f_{\theta}(x) = f_{\theta_0}(x) + \nabla_{\theta} f_{\theta_0}(x) \cdot (\theta - \theta_0) \\ + R_2(\theta; x)$$

$\exists \text{ se } (0,1)$

$$R_2(\theta; x) = \frac{1}{2} (\theta - \theta_0)^T$$

$$\left[\nabla_{\theta}^2 f_{\theta_0 + s(\theta - \theta_0)}(x) \right] (\theta - \theta_0)$$

harry training says

if a regime parameter

$$\|\theta(t) - \theta_0\| \rightarrow 0$$

and

$$f_{\theta}^{\text{true}}(x) := f_{\theta_0}(x) + \nabla_{\theta} f_{\theta_0}(x) \cdot (\theta - \theta_0)$$

becomes uniformly accurate
during training.

OK so this is the
key scaling mechanism

$$f_{\theta}(x) := \mathcal{L} \tilde{f}_{\theta}(x)$$

Define $\delta\theta := \theta - \theta_0$

$$\tilde{f}_\theta(x) = \tilde{f}_{\theta_0}(x) + J_{\theta_0}(x)\delta\theta + \dots$$

$$J_{\theta_0}(x) := \nabla_\theta \tilde{f}_{\theta_0}(x)^T \in \mathbb{R}^{1 \times p}$$

$$f_\theta(x) = \tilde{f}_{\theta_0}(x) + J_{\theta_0}(x)\delta\theta$$

Evaluating

$$\vec{f}(\theta) = \vec{\tilde{f}_0} + J\delta\theta + \dots$$

$$\mathcal{L}(\theta) = \frac{1}{2} \|\vec{\tilde{f}_0} + J\delta\theta + \dots - y\|_2^2$$

near a critical point

$$\Delta f_0 = \vec{y} - \Delta \tilde{f}_0 + \dots$$

$$\rightarrow \delta\theta = \frac{1}{2} J^+ (\vec{y} - \Delta \tilde{f}_0)$$

$$\|\delta\theta\| = O(1/\alpha)$$

↑ Assume small initialization!

And also ~~But~~ the tangent kernel

$$\Theta_0(x, x') = \Delta^2 \langle \nabla_0 \tilde{f}_0(x), \nabla_0 \tilde{f}_0(x') \rangle$$

If we assume (uniformly bounded Hessian)

$$\|\nabla_0 \tilde{f}_0(x) - \nabla_0 \tilde{f}_0(x')\| \leq H \|\theta - \theta'\|$$

Using $\| \theta - \theta_0 \| = O(1/k)$

~~✗~~ $\| \nabla_{\theta} \tilde{f}_{\theta}^{(k)} - \nabla_{\theta} \tilde{f}_{\theta_0}^{(k)} \| = O(k)$
 Bring it here

$$\Theta_0 - \Theta_{\theta_0} = \lambda^2 \left(\langle \nabla f(x), \nabla_{\theta} f(x') \rangle - \langle \nabla_{\theta} f(x), \nabla_{\theta} f(x') \rangle \right)$$

$$= \lambda^2 \left(\langle \tilde{f}_{\theta}^{(k)}(x) - \nabla \tilde{f}_{\theta_0}^{(k)}(x), \nabla \tilde{f}_{\theta}^{(k)}(x') \rangle + \langle \tilde{f}_{\theta_0}^{(k)}(x), \nabla \tilde{f}_{\theta}^{(k)}(x') - \nabla \tilde{f}_{\theta_0}^{(k)}(x') \rangle \right)$$

$$\| \Theta_{\theta} - \Theta_{\theta_0} \| = \lambda^2 O\left(\frac{1}{k}\right) = O(k)$$

$$\frac{\|\Theta - \Theta_0\|}{\|\Theta_0\|} = O(1/L)$$

nearly constant

Lazy training is a

scaling regime, not

a structural property of NN's.

Large L means:

- tiny movement in parameters
- produces larger changes in f

informal calculus

$$\tilde{f}_{\theta_0 + \delta\theta}(x) = \tilde{f}_{\theta_0}(x) + \nabla_{\theta} \tilde{f}_{\theta_0}(x) \cdot \delta\theta + \dots$$

multiply through by α

$$\delta f = \alpha \nabla_{\theta} \tilde{f}_{\theta_0} \cdot \delta\theta + \dots$$

$$|\delta f| \sim \alpha \|\delta\theta\|$$