

1/28/2026 AMPD UP (NTK)

I. NTK

A. Data.

1. ~~fast~~ consider fully-connected NN w/ layers $0, \dots, L$
 - a. layers have $n_0, \dots, n_L$ neurons
 - b. $\sigma: \mathbb{R} \rightarrow \mathbb{R}$, odd and derivative for nonlinearity
2. the realization f_θ for the NN maps $\theta \mapsto f_\theta$ (our NN)
 - a. dlm. $\theta = \sum_{l=0}^{L-1} (W^{(l+1)} \theta^{(l)} + b^{(l+1)})$ bias

3. parameters

- a. connection matrices $W^{(l)} \in \mathbb{R}^{n_l \times n_{l-1}}$
- b. bias vectors $b^{(l)} \in \mathbb{R}^{n_l}$
- c. initialize iid from $\mathcal{N}(0, 1)$

literally just weighted ℓ^2 norm $\rightarrow$ 4. consider p^{in} dist'n on $\mathbb{R}^{n_0}$ and then define seminorm $\|\cdot\|_{p^{in}}$ by $\langle f, g \rangle_{p^{in}} = \mathbb{E}_{x \sim p^{in}} (f(x)^T g(x))$

- a. In particular, take p^{in} to be empirical dist'n $p^{in} = \frac{1}{N} \sum_{i=1}^N \delta_{x_i}$

4. In particular,

$$f_\theta(x) = \tilde{z}^{(L)}(x; \theta)$$

a. $\tilde{z}^{(l)}(\cdot; \theta): \mathbb{R}^{n_0} \rightarrow \mathbb{R}^{n_l}$ - preactivations

b. $z^{(l)}(\cdot; \theta): \mathbb{R}^{n_l} \rightarrow \mathbb{R}^{n_l}$ - activations

c. define recursively:

$$i. \quad \tilde{z}^{(0)}(x; \theta) = x$$

$$ii. \quad \tilde{z}^{(l+1)}(x; \theta) = \frac{1}{\sqrt{n_l}} W^{(l+1)} \tilde{z}^{(l)}(x; \theta) + \beta b^{(l+1)}$$

$$iii. \quad z^{(l)}(x; \theta) = \sigma(\tilde{z}^{(l)}(x; \theta))$$

d. here, σ is applied entrywise

e. β is a hyperparameter

5. note by dividing by $\sqrt{n_l}$ helps when take width large.

a. and β helps $b^{(l)}$ from dominating when n_l is large if β is small enough

B. Kernel gradient

1. a multi-dimensional kernel is a fun $\mathbb{R}^{n_0} \times \mathbb{R}^{n_0} \rightarrow \mathbb{R}^{n_L \times n_L}$

review:
why do they divide by $\sqrt{n_l}$?

13) such that for any $x, x' \in \mathbb{R}^n$

$$K(x, x') = K(x', x)^T$$

a. equivalently, K is a symmetric tensor in $\mathcal{F} \otimes \mathcal{F}$
 i. think of a basis for $\mathcal{F}$ as $\left\{ \begin{pmatrix} 0 \\ \delta_i \\ 0 \end{pmatrix} \right\}$, then

basis
 el't of $\mathcal{F} \otimes \mathcal{F}$ looks like $\begin{pmatrix} 0 & \dots & 0 \\ \delta_i & & \delta_j \\ 0 & & 0 \end{pmatrix}$

2. given kernel, define bilinear map on $\mathcal{F}$ by
 $\langle f, g \rangle_K = \sum_{x, x'} f(x) K(x, x') g(x')$

independent

review LCTVS stuff

3. Def'n: K is pos. def. wrt $\|\cdot\|_{\text{pin}}$ if
 $\|f\|_{\text{pin}} > 0 \Rightarrow \|f\|_K > 0$

4. $\mathcal{F}^* = \text{dual of } \mathcal{F} \text{ wrt pin}$

Since have inner product, can use Riesz repr?

a. $\mathcal{B}$ set of linear forms $\mu: \mathcal{F} \rightarrow \mathbb{R}$ of the form

$\mu = \langle d, \cdot \rangle_{\text{pin}}$ for some $d \in \mathcal{F}$

5. let $K_{ij}(x_i, x_j)$ denote "partial application of the kernel"

6. define map $\Phi_K: \mathcal{F}^* \rightarrow \mathcal{F}$ by

$$\Phi_K(\mu)(x) = \mu(K_{ij}(x_i, \cdot)) = \langle d, K_{ij}(x_i, \cdot) \rangle_{\text{pin}}$$

where $\mu = \langle d, \cdot \rangle_{\text{pin}}$

a. write $f_\mu = \Phi_K(\mu)$

7. Now, cost functional C depends only on our data points $x_1, \dots, x_n$

8. we can compute variational derivative of C —

is an el't of $\mathcal{F}^*$

a. by definition if $f_0 \in \mathcal{F}$ is denoted $\mathcal{D}_f^{in} C|_{f_0}$

b. if we write $d|_{f_0}$ for el't of $\mathcal{F}$ s.t.

$$\mathcal{D}_f^{in} C|_{f_0} = \langle d|_{f_0}, \cdot \rangle_{\text{pin}}$$

9. define kernel gradient as

$$\nabla_K C|_{f_0} = \Phi_K(\mathcal{D}_f^{in} C|_{f_0})$$

a. we have

$$\nabla_K C|_{f_0}(x) = \frac{1}{n} \sum_{j=1}^n K(x, x_j) d|_{f_0}(x_j)$$

~~to make this is mean of a "family" generated by a family of functions...~~

10. Hence, could think of kernel gradient as trying to find approximately $\frac{\partial C}{\partial \theta}$ on space w/ both measure than pin by using some kernel to add up its denominator w/ these measures

11. a. fcn f follows the kernel gradient descent wrt K if it satisfies

$$\partial_t f(t) = -\nabla_K C(f(t))$$

a. during such descent, cost evolves per

$$\begin{aligned} \partial_t C(f(t)) &= -\langle d(f(t)), \nabla_K C(f(t)) \rangle_{\text{pin}} \\ &= -\|d(f(t))\|_K^2 \end{aligned}$$

12. Thus, if C is convex & bdd from below & K is pos. def. wrt $\|\cdot\|_{\text{pin}}$, then $f(t)$ will converge to a global min as $t \rightarrow \infty$.

13. ex.: Random. Run approximation

a. Pick some fcn $f^{(p)} \in \mathcal{F}$, $p=1, \dots, P$

b. Use these to define a parametrization $F^{\text{lin}}: \Theta \rightarrow \mathcal{F}$

$$\partial_\theta F^{\text{lin}}(\theta) = \frac{1}{P} f^{(p)}$$

given by

$$F^{\text{lin}}(\theta) = \frac{1}{P} \sum_{p=1}^P \theta_p f^{(p)}$$

c. Next, say params evolve through grad. descent:

$$\begin{aligned} \partial_t f(t) &= -\partial_\theta (C \circ F^{\text{lin}})(\theta(t)) \\ &= -\sum_{i=1}^n \partial_i C|_{f(t)} \partial_\theta F^{\text{lin}}(\theta(t)) \end{aligned}$$

$$= -\frac{1}{P} \sum_{p=1}^P \partial_i C|_{f(t)} f^{(p)} = -\frac{1}{P} \langle d|_{f(t)}^{\text{lin}}, f^{(p)} \rangle_{\text{pin}}$$

c. so $f_{\text{OCH}}^{\text{lin}} = F^{\text{lin}}(\theta(t))$ evolves per

$$\begin{aligned} \partial_t f_{\text{OCH}}^{\text{lin}} &= \frac{1}{P} \sum_{p=1}^P \partial_t \theta_p(t) f^{(p)} \\ &= -\frac{1}{P} \sum_{p=1}^P \langle d|_{f(t)}^{\text{lin}}, f^{(p)} \rangle_{\text{pin}} f^{(p)} \end{aligned}$$

f. but for $\tilde{C}$ for

$$\tilde{C} = \frac{1}{P} \sum_{p=1}^P \partial_i C|_{f(t)} f^{(p)}$$

$$F_{\text{lin}}^{\text{lin}}(x, \omega) = \frac{1}{P} \sum_{p=1}^P \partial_i C|_{f(t)} f^{(p)}(x, \omega)$$

c. Neural Tangent Kernel (NTK)

1. For real set up, we have

$$\begin{aligned} \frac{1}{2} C(a, \partial \theta_p &= -\partial_p (C \circ F^{(L)}) (\theta(k)) \\ &= -\partial_p^T C|_{F^{(L)}(\theta(k))} \partial_p F^{(L)} (\theta(k)) \\ &= -\langle d|_{F^{(L)}(\theta(k))}, \partial_p F^{(L)} (\theta(k)) \rangle_{p^m} \end{aligned}$$

6. so

$$\begin{aligned} \partial f_{\text{out}} &= \sum_{p=1}^P \partial \theta_p \partial_p F^{(L)} (\theta(k)) \\ &= \sum_{p=1}^P -\langle d|_{F^{(L)}(\theta(k))}, \partial_p F^{(L)} (\theta(k)) \rangle_{p^m} \partial_p F^{(L)} (\theta(k)) \\ &= \sum_{p=1}^P \sum_{j=1}^N (d|_{F^{(L)}(\theta(k))} (x_j)^T \partial_p F^{(L)} (\theta(k)) (x_j) \partial_p F^{(L)} (\theta(k)) \\ &= \sum_{p=1}^P \sum_{j=1}^N \partial_p F^{(L)} (\theta(k)) \partial_p F^{(L)} (\theta(k)) (x_j)^T d|_{F^{(L)}(\theta(k))} (x_j) \end{aligned}$$

THIS IS
the
NTIC

c. Hence, we set

$$\Theta^{(L)} (\theta) = \sum_{p=1}^P \partial_p F^{(L)} (\theta) \otimes \partial_p F^{(L)} (\theta)$$

d. then $\sum_{j=1}^N f_{j,k} = \Theta^{(L)} (\theta) (x_1, x_j)$, so

$$\partial f_{\text{out}} = -\nabla_{\Theta^{(L)}} C|_{f_{\text{out}}}$$

2. Proposition: For a network of depth L at initialization w/ Lipschitz non linearity σ , in the limit $n_1, \dots, n_L \rightarrow \infty$, the output functions $f_{j,k} (k=1, \dots, n_L)$ tend (in law) to iid centered Gaussian processes of covariance $\Sigma^{(L)}$ defined by

a. $\Sigma^{(L)} (x, x') = \frac{1}{n_0} x^T x' + \beta^2$

b. $\Sigma^{(L+1)} (x, x') = \mathbb{E}_{f \sim \pi(\theta, \Sigma^{(L)})} (\sigma(f(x)) \sigma(f(x'))) + \beta^2$,
taking the expectation wrt a centered Gaussian process f of covariance $\Sigma^{(L)}$.

3. so:

a. $f_{j,k}$ is a Gaussian

b. but! since θ is randomly initialized, for any fixed input x , the value $f_{j,k}(x)$ may be seen as a rv

c. and we're just saying what the law of this rv is in limit as width $\rightarrow \infty$.

4. prove B by induction

a. In base step, we have covariance Σ

$$\begin{aligned}\Sigma^{(1)}(x, x') &= \mathbb{E}_{\theta \sim q(\theta)} (f_\theta(x) f_\theta(x')^T) \\ &= \mathbb{E}_{\theta \sim \mathcal{N}(0, I)} \left(\frac{1}{\sqrt{n_0}} (W^{(0)} x + \beta b^{(0)}) \frac{1}{\sqrt{n_0}} (W^{(0)} x' + \beta b^{(0)})^T \right) \\ &= \frac{1}{n_0} \mathbb{E} \left(W^{(0)} x (W^{(0)} x')^T + W^{(0)} x \beta b^{(0)T} \right. \\ &\quad \left. + \beta b^{(0)} (W^{(0)} x')^T + \beta^2 b^{(0)} b^{(0)T} \right)\end{aligned}$$

use iid $\eta^{(0)}$

$$\text{assumption} \rightarrow \frac{1}{n_0} x^T x' I + 0 + 0 + \beta^2 I$$

5.4.4 a. so f_θ are independent (diagonal matrix) and all have the indicated covariance, as desired.

b. in inductive step, say claim holds for L layers

c. so as $n_1, \dots, n_{L-1} \rightarrow \infty$, $\tilde{\Sigma}_i^{(L)}$ tends to iid Gaussians w/covariance $\Sigma^{(L)}$

d. so we must consider

$$f_{\theta, i} = \frac{1}{\sqrt{n_L}} W_i^{(L)} \alpha^{(L)} + \beta b_i^{(L)}$$

e. exactly like above, these have covariance

$$\Sigma^{(L+1)}(x, x') = \frac{1}{n_L} \alpha^{(L)}(x)^T \alpha^{(L)}(x') + \beta^2$$

f. so LLN: in limit, set

$$\Sigma^{(L+1)}(x, x') = \mathbb{E}_{\theta \sim \mathcal{N}(0, \Sigma^{(L)})} (\sigma(f(x)) \sigma(f(x'))) + \beta^2$$

5. Thm: For a network of depth L at initialization w/lipschitz nonlinearity σ , in the limit as the widths $n_1, \dots, n_{L-1} \rightarrow \infty$, the NTK $\Theta^{(L)}$ converges in probability to a deterministic limiting kernel

Θ is

unnecessary

here

$$\Theta^{(L)} \rightarrow \Theta_{\infty}^{(L)} \oplus \text{Id}_{\text{out}}$$

where scalar kernel $\Theta_{\infty}^{(L)} : \mathbb{R}^{n_0} \times \mathbb{R}^{n_0} \rightarrow \mathbb{R}$ is

let's b₂

$$a. \Theta_{\infty}^{(L)}(x, x') = \Sigma^{(L)}(x, x')$$

$$b. \Theta_{\infty}^{(L+1)}(x, x') = \Theta_{\infty}^{(L)}(x, x') \dot{\Sigma}^{(L+1)}(x, x') + \Sigma^{(L+1)}(x, x')$$

$$\text{for } \dot{\Sigma}^{(L+1)}(x, x') = \mathbb{E}_{\theta \sim \mathcal{N}(0, \Sigma^{(L)})} (\dot{\sigma}(f(x)) \dot{\sigma}(f(x')))$$

i. with expectation wrt centered Gaussian process of covariance $\Sigma^{(L)}$

6 proof is by induction

as for $L=1$, use

$$(\partial_{w_{ij}} F)_k = \delta_{ik} x_j \frac{1}{\sqrt{n_0}}$$

i. so

$$\partial_{w_{ij}} F (\partial_{w_{ij}} F)^T = \frac{1}{n_0} \delta_{ik} x_j x'_j$$

ii. so

$$\Theta^{(1)}(\theta)(x, x') = \frac{1}{n_0} x^T x' I_{n_0} + \beta^2 = \Sigma^{(1)}(x, x')$$

1.6.6. for induction, say holds for L layers

as $n_1, \dots, n_{L-1} \rightarrow \infty$, $\tilde{\Sigma}_i^{(L)}$ goes to iid

Gaussian w/covariance $\Sigma^{(L)}$ per prev. proposition

d. Look at $\tilde{\Theta}_p$ belongs to first L layers, then

$$\partial_{\tilde{\Theta}_p} F^{(L)} = \frac{1}{n_L} W^{(L+1)} \partial \phi(\tilde{z}^{(L)}) \partial_{\tilde{\Theta}_p} \tilde{z}^{(L)}$$

Also for these parameters, get

$$\mathbb{E} \partial_{\tilde{\Theta}_p} F^{(L)}(w) (\partial_{\tilde{\Theta}_p} F^{(L)}(w))^T$$

$$= \mathbb{E} \frac{1}{n_L} W^{(L+1)} \partial \phi(\tilde{z}^{(L)}) \partial_{\tilde{\Theta}_p} \tilde{z}^{(L)} \partial_{\tilde{\Theta}_p} \tilde{z}^{(L)T} \partial \phi(\tilde{z}^{(L)T}) W^{(L+1)T}$$

$$\text{I.H.} \rightarrow n_1, \dots, n_{L-1} \rightarrow \infty \rightarrow \frac{1}{n_L} \Theta_{\tilde{\Theta}_p}^{(L)} W^{(L+1)} \partial \phi(\tilde{z}^{(L)}) \partial \phi(\tilde{z}^{(L)T}) W^{(L+1)T}$$

$$\text{LW, induce} \rightarrow n_L \rightarrow \infty \rightarrow \Theta_{\tilde{\Theta}_p}^{(L)} \text{Id} \mathbb{E} [\partial \phi(w) \partial \phi(w)^T]$$

e. next, for other params, just use

$$\partial_{w_j} F^{(L)} = \frac{1}{n_L} \phi(\tilde{z}_j^{(L)}) e_i$$

$$\text{f. so } \mathbb{E} \partial_{w_j} F^{(L)} (\partial_{w_j} F^{(L)})^T = \frac{1}{n_L} \mathbb{E} \phi(\tilde{z}_j^{(L)}) \phi(\tilde{z}_j^{(L)}) \text{Id}$$

g. use proposition, don't forget β_i^2 you're good.

V. Training

just bound. 1. general defin of training: update param's per

$$\partial_{\tilde{\Theta}_p}(t) = \langle \partial_{\tilde{\Theta}_p} F^{(L)}(t), d_t \rangle_{\text{pin}}$$

what does

$\rightarrow$ where $\int_0^T \text{Idellpin} dt$ is "stochastically bounded"

"stochastically

" width goes to ∞

bounded"

b. for gradient descent, take

$$d_t = -d|_{f(t)}$$

2. Thm: Assume ϕ Lipschitz, 2nd db nonlin

non unbounded 2nd derivate. For any T s.t.

$\int_0^T \text{Idellpin} dt$ stays stochastically bounded, or

or

$$\partial_{\tilde{\Theta}_p} f(t, w)$$

$$= \langle d_t, \Theta_{\tilde{\Theta}_p}^{(L)}(t) \text{Id} \rangle$$

$n_1, \dots, n_{L-1} \rightarrow \infty$, we have, uniformly for $t \in [0, T]$,

$$\Theta_{\tilde{\Theta}_p}^{(L)}(t) \rightarrow \Theta_{\tilde{\Theta}_p}^{(L)} \text{Id}_{n_L}$$

a. As a consequence, in long limit, the dynamics of f_θ is described by the diff eq.

$$\partial_t f_\theta(t) = \Theta_{\tilde{\Theta}_p}^{(L)} \text{Id}_{n_L} (\langle d_t, \cdot \rangle_{\text{pin}})$$

3. proof β by induction

a. for $L=1$, NTK doesn't depend on param's, so it

by reverse triangle

$$\|g(t+h)\|_4 = \lim_{h \rightarrow 0}$$

$$\frac{\|g(t+h)\|_4 - \|g(t)\|_4}{h} \leq \lim_{h \rightarrow 0}$$

$$\frac{\|g(t+h)\|_4 - \|g(t)\|_4}{h} = \|g'(t)\|_4$$

constant

to 36. For other L , need a lemma

4. lemma: In setting of thm 2, for a network of depth $L+1$, for any $t \in [0, 1]$, we have convergence in probability

does not apply to $W^{(0)}$

$$\lim_{n \rightarrow \infty} \dots \lim_{n \rightarrow \infty} \sup_{t \in [0, 1]} \left\| \frac{1}{\sqrt{n}} (W^{(0)}(t) - W^{(0)}(0)) \right\|_q = 0$$

4. proof

a. we have

$$d_t W^{(l)} = \frac{1}{\sqrt{n}} \langle \alpha^{(l)}; g_t^{(l)} \rangle_{p, n}$$

$$d_t \Theta^{(l)} = \Theta^{(l)} \Theta^{(l+1)} \left(\langle d_t \alpha^{(l+1)}; \cdot \rangle_{p, n} \right)$$

for

$$d_t^{(l)} = \begin{cases} d_t & \text{if } l = L+1 \\ \Theta^{(l)} \frac{1}{\sqrt{n}} (W^{(l+1)})^T d_t^{(l+1)} & \text{if } l \leq L \end{cases}$$

why p, n listed of Euclidean norm?

$$b. \text{ set } w^{(l)}(t) = \left\| \frac{1}{\sqrt{n}} W^{(l)}(t) \right\|_{\text{op}}, \quad a^{(l)}(t) = \left\| \frac{1}{\sqrt{n}} \alpha^{(l)}(t) \right\|_{p, n}$$

c. then

$$\|d_t^{(l)}\|_{p, n} \leq C w^{(l)}(t) \|d_t^{(l+1)}\|_{p, n}$$

$$\text{Lipshitz constant} \leq \dots \leq C^{L+1} e^{-\frac{1}{2} \sum_{l=1}^L w^{(l)}(t)} \|d_t^{(L+1)}\|_{p, n}$$

d. Note sub networks can be written

$$\Theta^{(l)} = \frac{1}{n_0} \alpha^{(l+1)T} \alpha^{(l)} \text{Id}_{n_l} + \beta^2 \text{Id}_{n_l}$$

$$\Theta^{(l+1)} = \frac{1}{n_l} W^{(l)} \Theta^{(l)} \Theta^{(l+1)} \Theta^{(l+1)T} W^{(l+1)T} + \frac{1}{n_l} \alpha^{(l+1)T} \alpha^{(l)} \text{Id}_{n_l} + \beta^2 \text{Id}_{n_l}$$

e. so set bounds

$$\|\Theta^{(l)}\|_{\text{op}} \leq (a^{(l+1)}(t))^2 + \beta^2$$

$$\|\Theta^{(l+1)}\|_{\text{op}} \leq a(w^{(l)}(t))^2 \|\Theta^{(l)}\|_{\text{op}} + (a^{(l+1)}(t))^2 + \beta^2$$

$$\leq \dots \leq P(a^{(1)}(t), \dots, a^{(L)}(t), w^{(1)}(t), \dots, w^{(L)}(t))$$

for P some poly.

P. now, define

$$\tilde{a}^{(l)}(t) = \left\| \frac{1}{\sqrt{n_l}} (\alpha^{(l)}(t) - \alpha^{(l)}(0)) \right\|_{p, n}$$

$$\tilde{w}^{(l)}(t) = \left\| \frac{1}{\sqrt{n_l}} (W^{(l)}(t) - W^{(l)}(0)) \right\|_{\text{op}}$$

$$A(t) = \sum_{l=1}^L a^{(l)}(t) + C \tilde{a}^{(L)}(t) + w^{(0)}(0) + w^{(0)}(t)$$

$$i. \text{ note } d_t^{(l)}(t), w^{(l)}(t) \leq A(t)$$

g. next, by eqns a, evaluate $d_t W, \tilde{\alpha}$, we get

does this include p, n norm?

Θ ?

Yes:

$$\|K\|_{\text{op}} = \max_{\|x\|_{p, n} = 1} \|Kx\|_k$$

$$= \max_{\|x\|_{p, n} = 1} \left\| \sum_{k=1}^K (Kx)_k e_k \right\|_k$$

$$\text{So } \partial_{\theta} \tilde{w}^{(c)}(t) \leq \frac{1}{\sqrt{n}} a^{(c)}(t) \|d_{\theta}^{(c)}(t)\|_{p,m}$$

$$\partial_{\theta} a^{(c)}(t) \leq \frac{1}{\sqrt{n}} \|\Theta^{(c)}(t)\|_{op} \|d_{\theta}^{(c)}(t)\|_{p,m}$$

h. therefore

$$\partial_{\theta} A(t) \leq \sum_{c=1}^C \frac{1}{\sqrt{n}} \|\Theta^{(c)}(t)\|_{op} \|d_{\theta}^{(c)}(t)\|_{p,m} + \frac{1}{\sqrt{n}} a^{(c)}(t) \|d_{\theta}^{(c)}(t)\|_{p,m}$$

$$\leq \sqrt{\frac{1}{nm}} \sum_{c=1}^C \|\Theta^{(c)}(t)\|_{op} \|d_{\theta}^{(c)}(t)\|_{p,m} + \frac{1}{\sqrt{n}} \sum_{c=1}^C a^{(c)}(t) \|d_{\theta}^{(c)}(t)\|_{p,m}$$

$$\leq \sqrt{\frac{1}{nm}} \sum_{c=1}^C \|\Theta^{(c)}(t)\|_{op} \|d_{\theta}^{(c)}(t)\|_{p,m} + \frac{1}{\sqrt{n}} \sum_{c=1}^C a^{(c)}(t) \|d_{\theta}^{(c)}(t)\|_{p,m}$$

for new poly's $\mathcal{I}, \mathcal{R}$

i. Now: use a Gronwall lemma

i. for the $A(t) = \sum_{c=1}^C a^{(c)}(t) + w^{(c)}(t)$ is stochastically bounded as $n \rightarrow \infty$, and so is $\partial_{\theta} A(t)$

ii. b.c. of our earlier part, we actually

get that $A(t)$ stays within bound in $A(t)$

in, in only get $\partial_{\theta} A(t) \rightarrow 0$ within of $A(t) \rightarrow A(t)$

is, roughly,

$$A(t) \leq A(t) \leq A(t) \exp(\dots)$$

$\rightarrow 0$ as $n_1, \dots, n_C \rightarrow \infty$

3. very uninformative: weights and activations approach constant as width $\rightarrow \infty$

a. "Lazy Training in Differentiable Programming" by Chizat et al.

b. "On the theory of large non-linear models: when and why the tangent kernel is constant" by Liu et al.

c. both say it's b.c. the NN approaches a linear NN

about and does Hessian analysis

E. least squares regression

$$1. C(f) = \frac{1}{2} \|f - f^*\|_{p,m}^2 \quad \text{for target } f^*$$

2. this is the cost