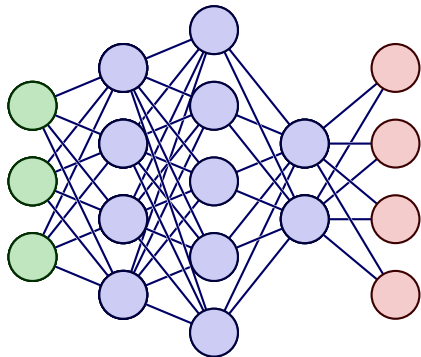


Training NNs for PDF Determinations

L Del Debbio

Higgs Centre for Theoretical Physics
University of Edinburgh



work in collaboration with

A Chiefa, R Kenway, T Giani, A Candido, G Petrillo, M Wilson

see also

- Idd et al (2022)
- Idd et al (2024)
- new, *preliminary* results

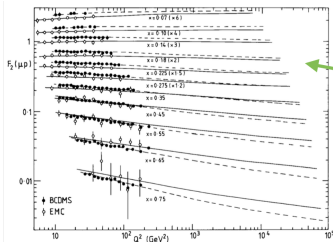
precision physics at the LHC

theoretical description of protons in terms of PDFs

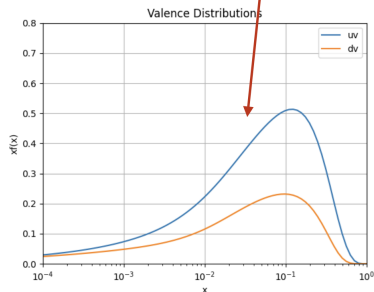
$$\text{e.g. for DIS, } T_I = \int dx C_I(x) f(x)$$

- requires precise and accurate determinations of the PDFs from data
- universality of PDFs: predictions for new processes
- typical example of an inverse problem
- NNPDF methodology: focus on the training process of NNs
- useful for LGTs: extraction of spectral densities/trivializing maps (normalizing flows)

inverse problem



$$y_I = \int dx C_I(x) f(x)$$

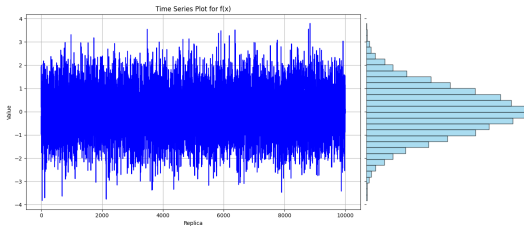


[NNPDF4.0]

parametrization of $f_i(x)$ for $x \in [0, 1]$: bias vs variance of the results

Bayesian framework

- promote f to a stochastic process, $f(x)$ are stochastic variables
- **choose** a prior distribution $p(f)$ - prior knowledge about f
- probability distributions are represented by ensembles of *replicas*

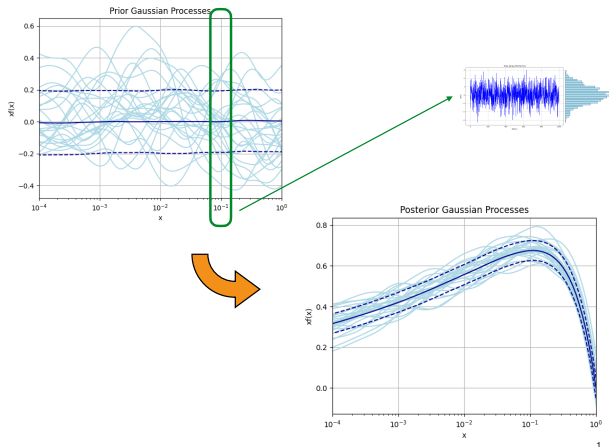


- observables are computed from averages over replicas

$$\bar{O} = E_p [O(f)] = \int df p(f) O(f) = \frac{1}{N_{\text{rep}}} \sum_{k=1}^{N_{\text{rep}}} O\left(f^{(k)}\right)$$

solving the inverse problem: posterior distribution

- use the data to infer a posterior probability $\tilde{p}(f)$

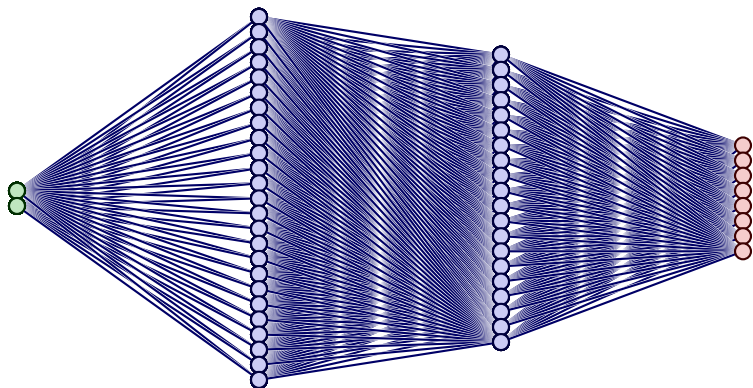


$$\bar{f}(x) = E_{\tilde{p}} [f(x)] , \quad \delta f(x) = E_{\tilde{p}} \left[(f(x) - \bar{f}(x))^2 \right]$$

NNPDF paradigm

- parametrize the unknown functions f using neural networks
- study a vector of values $f_\alpha = f(x_\alpha)$
- initialize N_{rep} NN replicas
- train each NN on a replica of the dataset
- training/validation split to determine the stopping condition
- the ensemble of trained NNs provides the posterior distribution

NNPDF parametrization - conventions



$$\rho_i^{(\ell)} = \rho(\phi_i^{(\ell)}), \quad \phi_i^{(\ell)} = \sum_{j=1}^{n_{\ell-1}} w_{ij}^{(\ell)} \rho_j^{(\ell-1)} + b_i^{(\ell)}$$

$$f(x) = \phi^{(L)}(x; \theta)$$

finite-dimensional, sufficiently large to be unbiased

NN prior distribution

parameters θ are initialized using a Glorot-Normal distribution

initialize weights and biases using Gaussians

$$\begin{aligned}\langle b_i^{(\ell)} \rangle &= 0, & \langle b_{i_1}^{(\ell)} b_{i_2}^{(\ell)} \rangle &= \delta_{i_1 i_2} C_b^{(\ell)} \\ \langle w_{ij}^{(\ell)} \rangle &= 0, & \langle w_{i_1 j_1}^{(\ell)} w_{i_2 j_2}^{(\ell)} \rangle &= \delta_{i_1 i_2} \delta_{j_1 j_2} \frac{C_w^{(\ell)}}{n_{\ell-1}}\end{aligned}$$

parameters/functions duality

$$p(\phi^{(\ell)}) = \int [dw p(w)] [db p(b)] \prod_{i,\alpha} \delta \left(\phi_{i\alpha}^{(\ell)} - \sum_j w_{ij}^{(\ell)} \rho \left(\phi_{j\alpha}^{(\ell-1)} \right) - b_i^{(\ell)} \right)$$

EFT approach

symmetry and $1/n$ counting yield the probability distribution for ϕ

$$\begin{aligned} p(\phi^{(\ell)}) &= \frac{1}{Z} \exp \left[-S(\phi^{(\ell)}) \right] \\ &= \frac{1}{Z} \exp \left[-\frac{1}{2} \gamma_{\alpha_1 \alpha_2}^{(\ell)} \phi_{\alpha_1}^{(\ell)} \cdot \phi_{\alpha_2}^{(\ell)} \right. \\ &\quad \left. - \frac{1}{8n_{\ell-1}} \gamma_{\alpha_1 \alpha_2, \alpha_3 \alpha_4}^{(\ell)} \phi_{\alpha_1}^{(\ell)} \cdot \phi_{\alpha_2}^{(\ell)} \phi_{\alpha_3}^{(\ell)} \cdot \phi_{\alpha_4}^{(\ell)} + O(1/n_{\ell-1}^2) \right] \end{aligned}$$

correlators can be computed using Feynman diagrams

$$\begin{aligned} \langle \phi_{i_1, \alpha_1}^{(\ell)} \phi_{i_2, \alpha_2}^{(\ell)} \rangle &= \delta_{i_1 i_2} K_{\alpha_1 \alpha_2}^{(\ell)} + O(1/n_{\ell-1}) \\ \langle \phi_{i_1, \alpha_1}^{(\ell)} \phi_{i_2, \alpha_2}^{(\ell)} \phi_{i_3, \alpha_3}^{(\ell)} \phi_{i_4, \alpha_4}^{(\ell)} \rangle_c &= O(1/n_{\ell-1}) \end{aligned}$$

ϕ are approximately Gaussian processes, corrections are $O(1/n)$

training

for all parametrizations and for a quadratic loss

$$\mathcal{L}_t = \frac{1}{2} (y - T[f_t])^T C_Y^{-1} \underbrace{(y - T[f_t])}_{\epsilon_t}$$

gradient descent

$$\begin{aligned} \frac{d}{dt} \theta_\mu &= -\nabla_\mu \mathcal{L} \\ \frac{d}{dt} f_t &= (\nabla_\mu f_t) \frac{d}{dt} \theta_\mu = \Theta_t \left(\frac{\partial T}{\partial f} \right)_t C_Y^{-1} \epsilon_t \end{aligned}$$

where

$$\Theta_t = (\nabla_\mu f_t)(\nabla_\mu f_t)^T$$

is the Neural Tangent Kernel (NTK)

linear data & NN parametrization

for linear data

$$T = (\text{FK})f \quad \Longrightarrow \quad \left(\frac{\partial T}{\partial f} \right) = (\text{FK})$$

for wide neural networks

$$\boxed{\Theta_t = \Theta + O(1/n) \quad (\text{lazy training})}$$

hence we obtain a linear equation for f_t

$$\begin{aligned} \frac{d}{dt} f_t &= \Theta(\text{FK})^T C_Y^{-1} (y - (\text{FK})f_t) \\ &= -\Theta M f_t + b \end{aligned}$$

where $M = (\text{FK})^T C_Y^{-1} (\text{FK})$

flat directions & NTK spectrum

eigenvalues of the NTK

$$\Theta z^{(k)} = \lambda^{(k)} z^{(k)}, \quad f_{t,k} = \left(z^{(k)}, f_t \right)$$

evolution equations in the NTK eigenbasis

$$\frac{d}{dt} f_{t,k} = \lambda^{(k)} \left(z^{(k)}, (\text{FK})^T C_Y^{-1} (y - (\text{FK}) f_t) \right)$$

NTK kernel: directions that do NOT evolve during training

$$f_{t,\parallel} = f_{0,\parallel}$$

no bias, but irreducible noise dictated by the prior

analytic solution

for each replica:

$$f_t = U(t)f_0 + V(t)Y$$

and therefore we have analytic expressions for

$$\bar{f}_t = \mathbb{E}[U(t)f_0] + \mathbb{E}[V(t)Y]$$

$$\text{Cov}[f_t, f_t^T] = C_t^{(00)} + C_t^{(0Y)} + C_t^{(YY)}$$

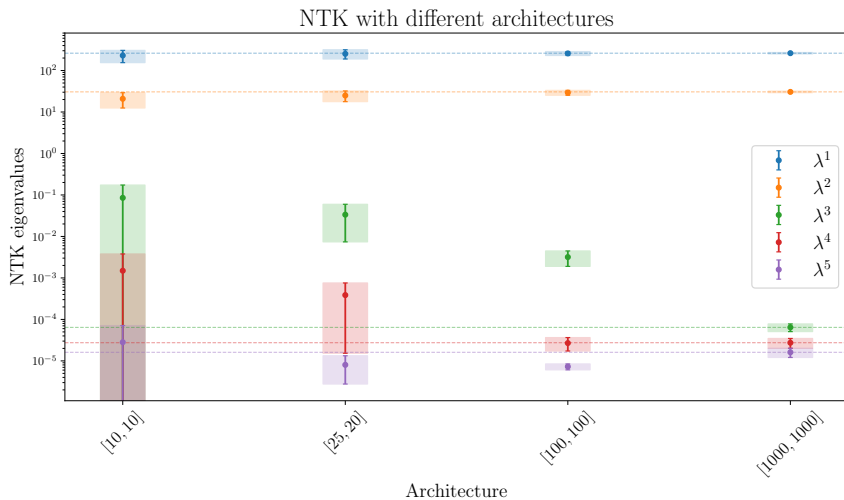
phenomenology

- for each replica $k = 1, \dots, N_{\text{rep}}$
- initialize a NN
- generate synthetic data using a known f^{in}

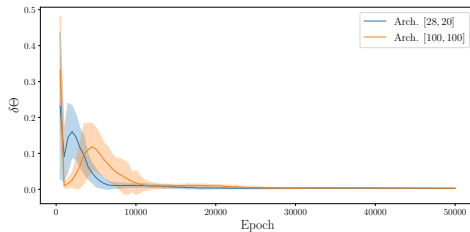
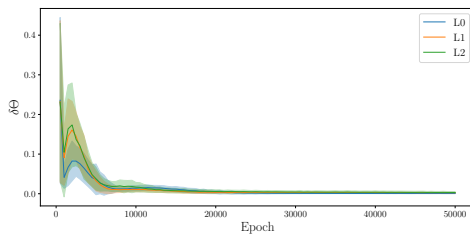
$$Y^{(k)} = \underbrace{(\text{FK})f^{\text{in}}}_{L_0} + \underbrace{\eta}_{L_1} + \underbrace{\xi^{(k)}}_{L_2}$$

- train numerically and compare with the analytic expressions

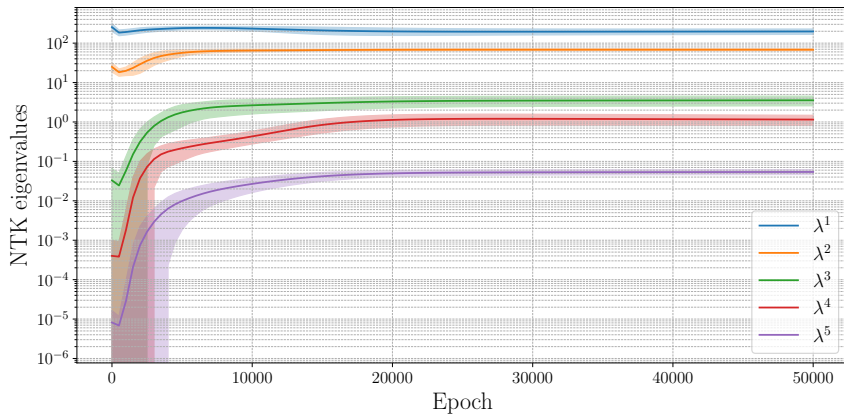
NTK at initialization



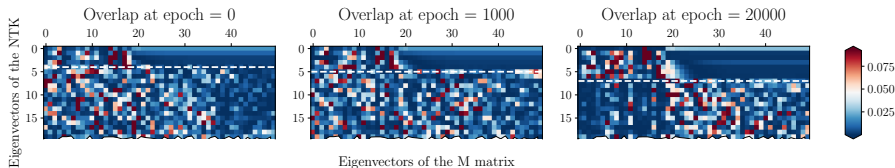
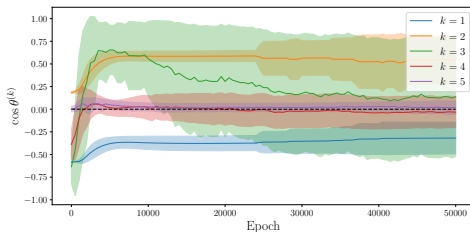
NTK during training



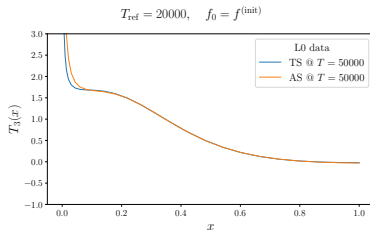
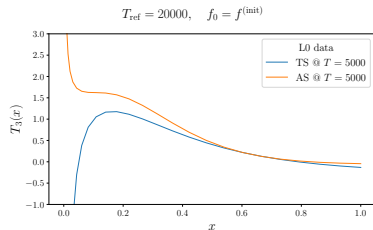
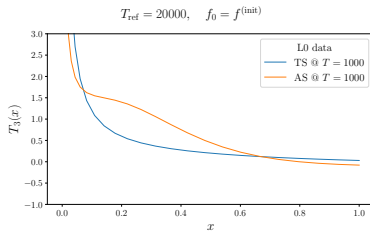
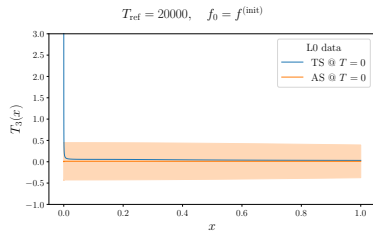
NTK alignment



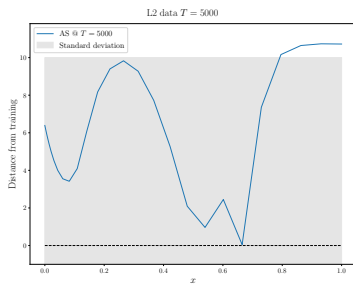
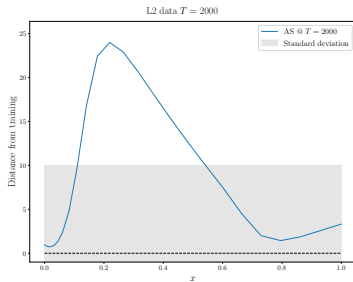
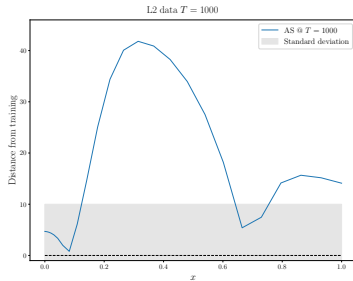
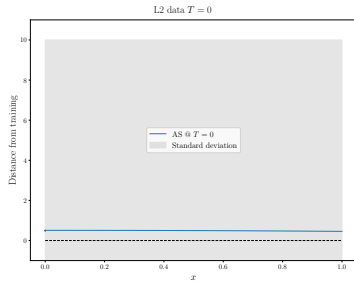
NTK alignment



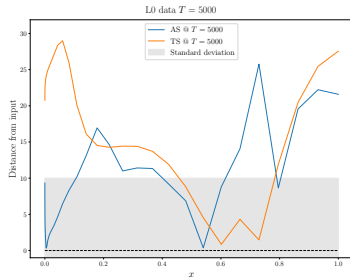
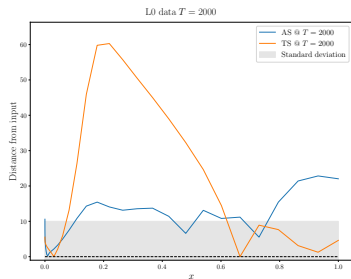
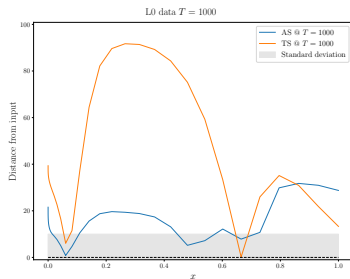
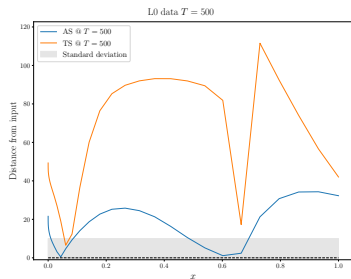
evolution with frozen NTK



distances

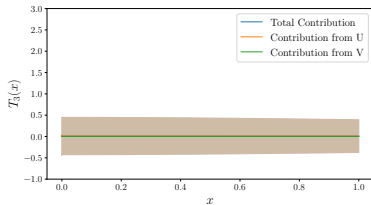


distances

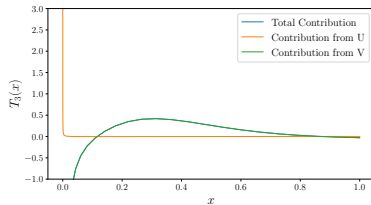


U and V contributions

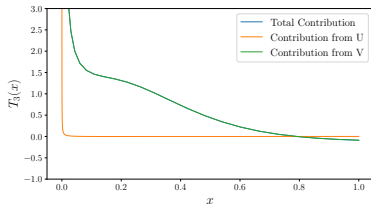
$$T_{\text{ref}} = 20000, \quad f_0 = f^{(\text{init})}, \quad T = 0$$



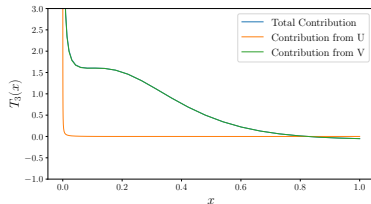
$$T_{\text{ref}} = 20000, \quad f_0 = f^{(\text{init})}, \quad T = 100$$



$$T_{\text{ref}} = 20000, \quad f_0 = f^{(\text{init})}, \quad T = 700$$



$$T_{\text{ref}} = 20000, \quad f_0 = f^{(\text{init})}, \quad T = 3000$$



outlook

- PDFs are a central ingredient for LHC analyses
- Bayesian approach is convenient to solve inverse problems
- all hypotheses are clearly spelled in the prior
- analytical control of the training process is key to quantify the bias/variance budget
- same methods can be applied to normalizing flows?