

# Machine Learning Approaches to Hard Exclusive Processes

Marija Čuić

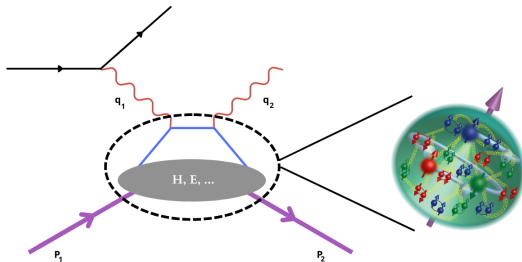
Irfu, CEA, Université Paris-Saclay, Aidas

Synergies between the EIC and the LHC

September 22 - 24, 2025, Krakow



# Nucleon structure



$$d\sigma \propto 4(1 - x_B) \left( |\mathcal{H}|^2 + |\tilde{\mathcal{H}}|^2 \right) - \dots$$

$$\downarrow$$

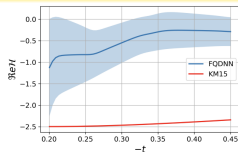
$$\mathcal{H}^A(\xi, \Delta^2, Q^2) = \underbrace{\int_{-1}^1 \frac{dx}{2\xi} A_T \left( x, \xi \middle| \alpha_s(\mu_R), \frac{Q^2}{\mu_F^2} \right)}_{\text{hard scale}} \underbrace{H^A(x, \eta, \Delta^2, \mu_F^2)}_{\text{soft scale}}$$

Inverse problem

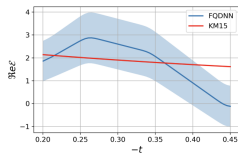
observables  $\rightarrow$  CFFs  $\rightarrow$  GPDs  $\Rightarrow$  ill-posed problem!

# Some recent CFF extractions

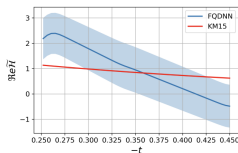
B. B. Le, D. Keller,  
arXiv:2504.15458,  
2025



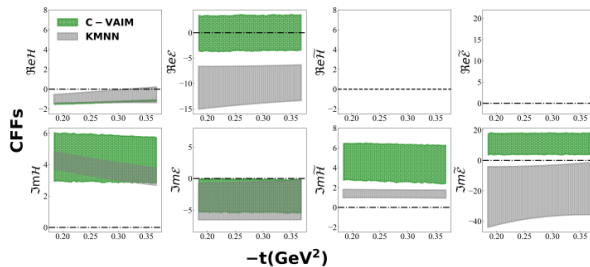
(a)  $\text{Re}H$  vs  $-t$  with  $x_B = 0.365$



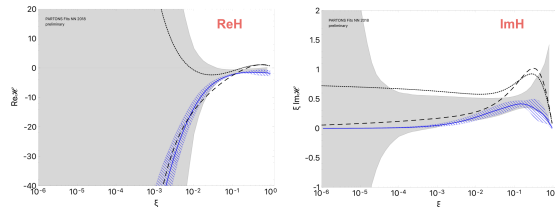
(b)  $\text{Re}E$  vs  $-t$  with  $x_B = 0.275$



(c)  $\text{Re}\tilde{H}$  vs  $-t$  with  $x_B = 0.305$

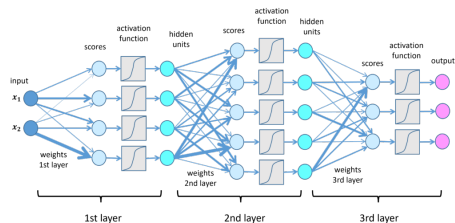


M. Almaen et al, 2024, vs MČ et al, PRL, 2020

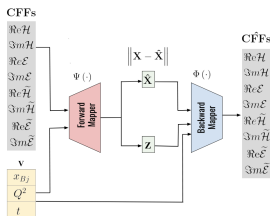


H. Moutarde et al, EPJ, 2019

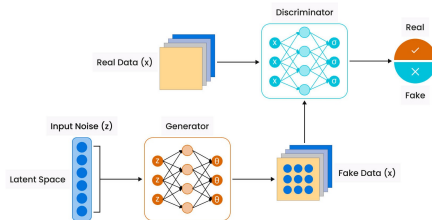
# Some machine learning methods on the market



Deep neural networks



Variational autoencoder inverse mapper

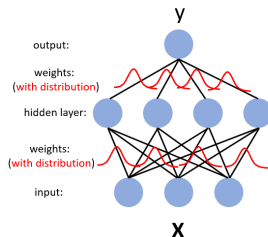


Generative adversarial network

and many more...



Symbolic regression



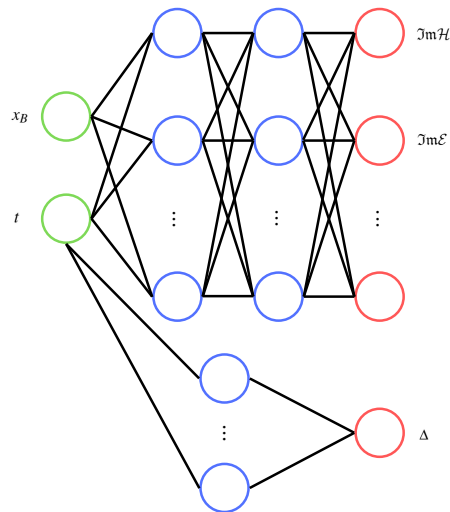
Bayesian neural networks

# Neural network architecture with dispersion relations

$$\Re \mathcal{H}(\xi, t) = \Delta(t) + \frac{1}{\pi} \text{P.V.} \int_0^1 dx \left( \frac{1}{\xi - x} - \frac{1}{\xi + x} \right) \Im \mathcal{H}(x, t)$$

We only model 4+1 functions instead of 8.

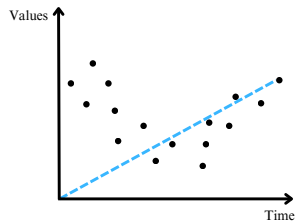
Loss function is reduced through gradient descent.



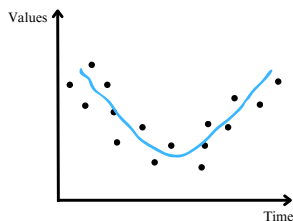
## Universal approximation theorem

Let  $C(K)$  be the space of continuous functions on a compact set  $K \subseteq \mathbb{R}^n$ . For any continuous function  $f \in C(K)$  and for any  $\varepsilon > 0$ , there exists a feedforward neural network  $\hat{f}$  with a single hidden layer such that:

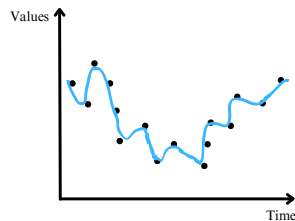
$$|f(x) - \hat{f}(x)| < \varepsilon \quad \text{for all } x \in K.$$



Underfitted



Good fit



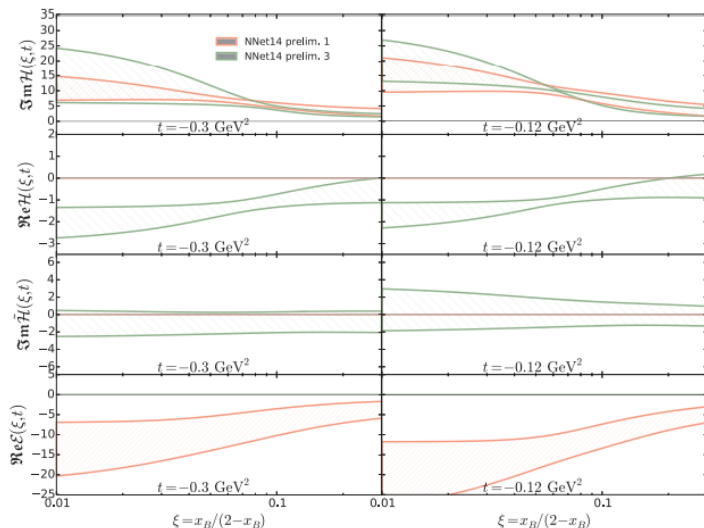
Overfitted

Model independent (mostly) and MC propagation of uncertainties.

# HERMES asymmetries

K. Kumerički, D. Müller, QCD  
Evol. 2014

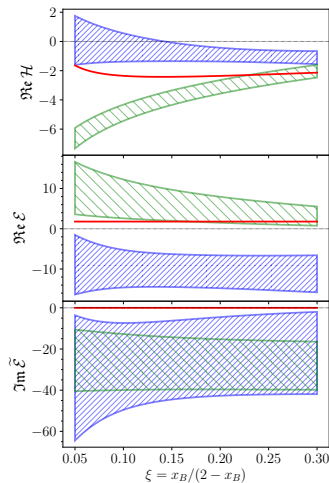
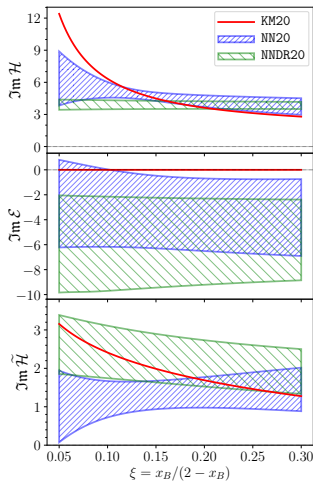
A good middle ground  
between global and local fits!



# JLab up to 6 GeV

MČ, K. Kumerički, A. Schäfer, PRL 2020

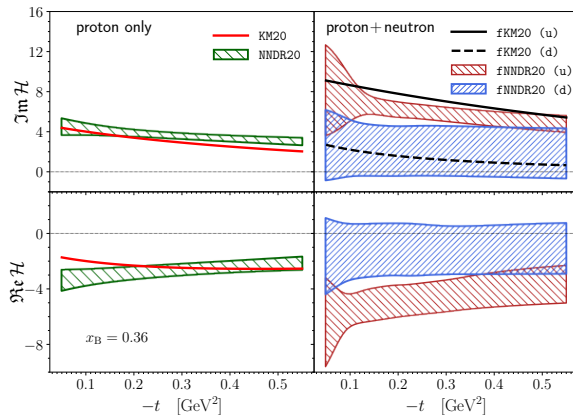
| Observable             | $n_{\text{pts}}$ | KM20 | NN20 | NNDR20 | fKM20 | fNNDR20 |
|------------------------|------------------|------|------|--------|-------|---------|
| # CFFs + $\Delta$ s    |                  | 3+1  | 6    | 4+1    | 5+2   | 8+2     |
| Total (harmonics)      | 277              | 1.3  | 1.6  | 1.7    | 1.7   | 1.8     |
| CLAS $A_{\text{LU}}$   | 162              | 0.9  | 1.0  | 1.1    | 1.2   | 1.3     |
| CLAS $A_{\text{UL}}$   | 160              | 1.5  | 1.7  | 1.8    | 1.8   | 2.0     |
| CLAS $A_{\text{LL}}$   | 166              | 1.3  | 3.9  | 0.8    | 1.1   | 1.6     |
| CLAS $d\sigma$         | 1014             | 1.1  | 1.0  | 1.2    | 1.2   | 1.1     |
| CLAS $\Delta\sigma$    | 1012             | 0.9  | 0.9  | 1.0    | 0.9   | 1.1     |
| Hall A $d\sigma$       | 240              | 1.2  | 1.9  | 1.7    | 0.9   | 1.3     |
| Hall A $\Delta\sigma$  | 358              | 0.7  | 0.8  | 0.8    | 0.7   | 0.7     |
| Hall A $d\sigma$       | 450              | 1.5  | 1.6  | 1.7    | 1.9   | 2.0     |
| Hall A $\Delta\sigma$  | 360              | 1.6  | 2.2  | 2.2    | 1.9   | 1.7     |
| Hall A $d\sigma_n$     | 96               |      |      |        | 1.2   | 0.9     |
| Total ( $\phi$ -space) | 4018             | 1.1  | 1.3  | 1.3    | 1.2   | 1.3     |



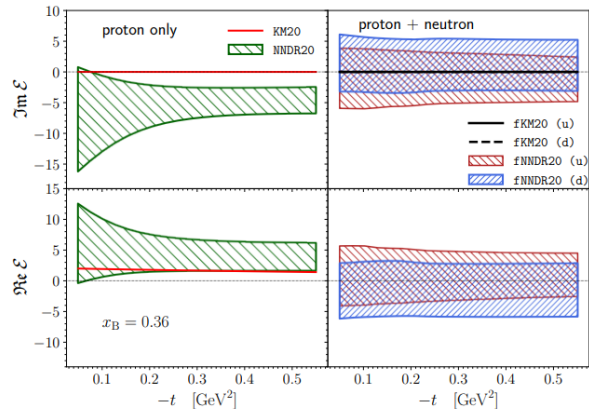


# Flavor separation with neutron DVCS

We separately model  $\mathcal{F}^u$  and  $\mathcal{F}^d$ .



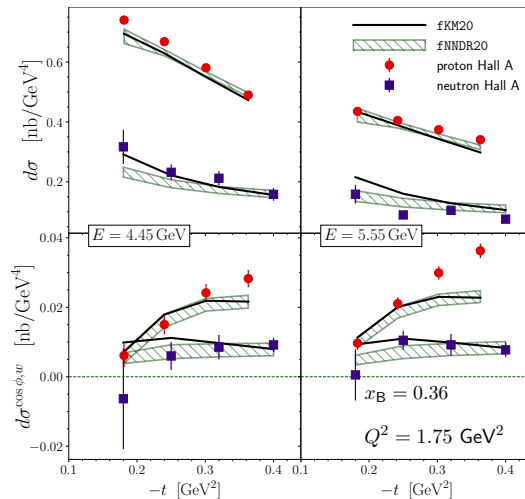
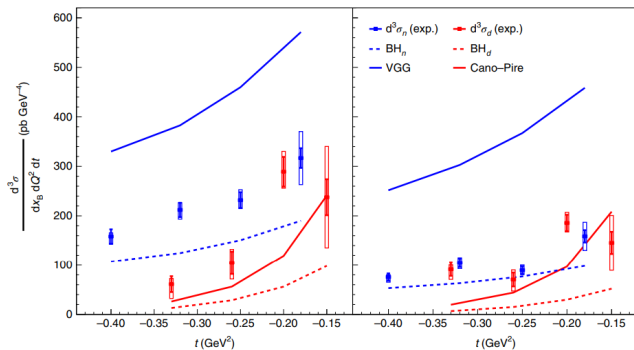
(a)  $\mathcal{H}$  flavor separation



(b)  $\mathcal{E}$  flavor separation

# Data representation

Benali et al. '20, DVCS off a deuterium target



# JLab 12 GeV upgrade

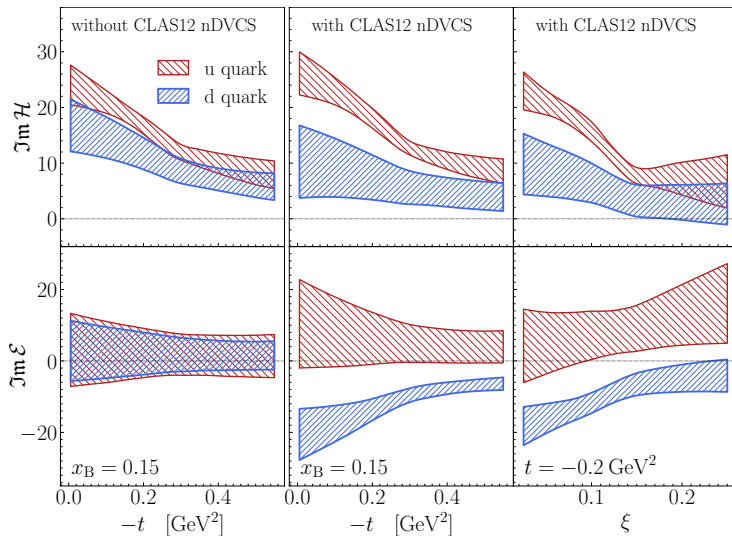
Fits from 2020 poorly represented the new BSA proton and neutron data. We performed new NN and NNDR fits on harmonics with new CLAS 2022 data (39 points with  $-t < 0.5 \text{ GeV}^2$ ) and previously available JLab data (257 points).

$$A_{LU} = \frac{d\sigma^{\uparrow} - d\sigma^{\downarrow}}{d\sigma^{\uparrow} + d\sigma^{\downarrow}} \propto \Im \left\{ F_1 \mathcal{H} + \xi (F_1 + F_2) \tilde{\mathcal{H}} - \frac{\Delta^2}{4M^2} F_2 \mathcal{E} \right\} \sin(\phi)$$

|       | 2020 | 2023 |
|-------|------|------|
| fNN   | 1.5  | 1.25 |
| fNNDR | 1.5  | >3   |

Table:  $\chi^2/N_{pts}$

# New flavor separation

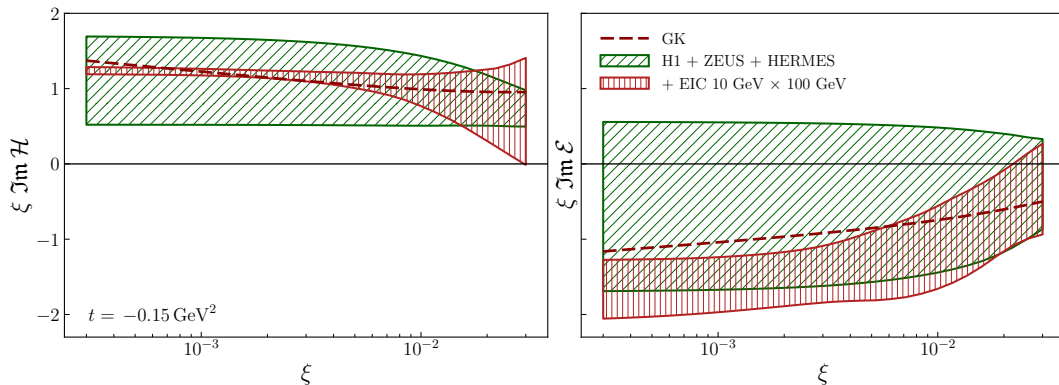


A. Hobart, S. Niccolai,  
MČ, K. Kumerički et al,  
PRL, 2024

We refitted only  
imaginary CFFs to  $A_{LU}$ ,  
 $A_{UL}$ ,  $X_{LU}$  and  $X_{UU}$ .  
Flavor separation of  
 $\Re \mathcal{H}$  is lost.

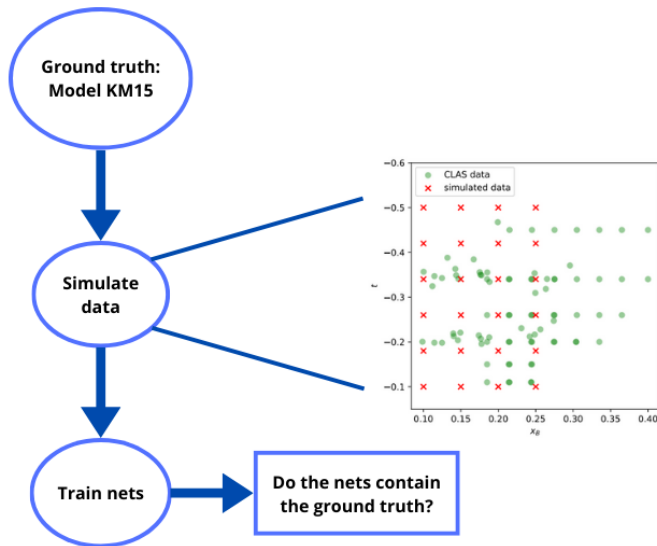
# Simulated EIC data

E. C. Aschenauer et al, PRD, 2025, extraction from simulated beam-spin asymmetry



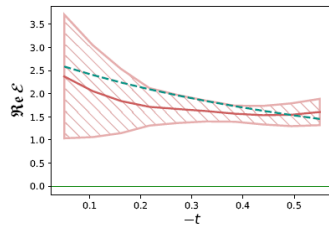
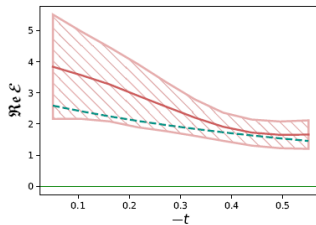
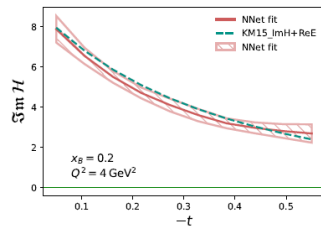
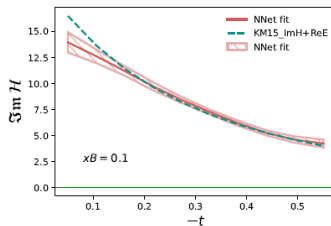
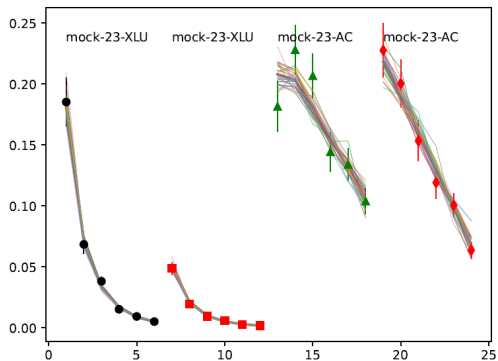
# How reliable are these results?

# Testing our method with closure tests



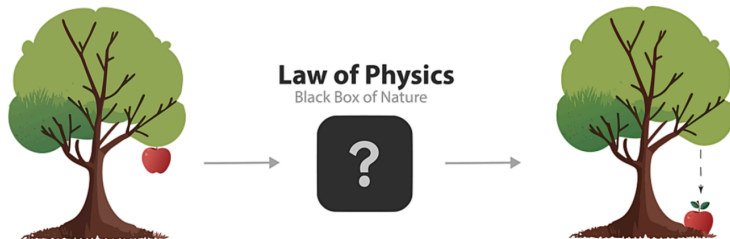
Which observables are required for CFF extraction?

# $\text{Im}\mathcal{H}(x_B, t)$ and $\text{Re}\mathcal{E}(x_B, t)$ from $X_{\text{LU}}$ and $A_C$



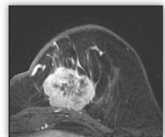


# Are neural networks black boxes? Why do they do what they do?



**Law of Physics**  
Black Box of Nature

Can we criticize the model?  
Can we interpret its inner workings? How? Why? Can we build it with interpretability in mind?



**Artificial Intelligence**  
Black Box Model



L. Sharkey et al, arXiv:2501.16496, 2025  
M. T. Ribeiro et al, proceeding, 2016

**All of physics is a black box!**

Figure: E. Marcus, J.  
Teuwen, Eur. J. Rad, 2024

# Outtakes

- Is the data too noisy for reliable extraction?
- Do we understand systematic uncertainties and experimental methods well enough for precise extraction?
- Can we benchmark?
- Are more sophisticated ML methods really necessary?
- Can we get reliable results in kinematic gaps?