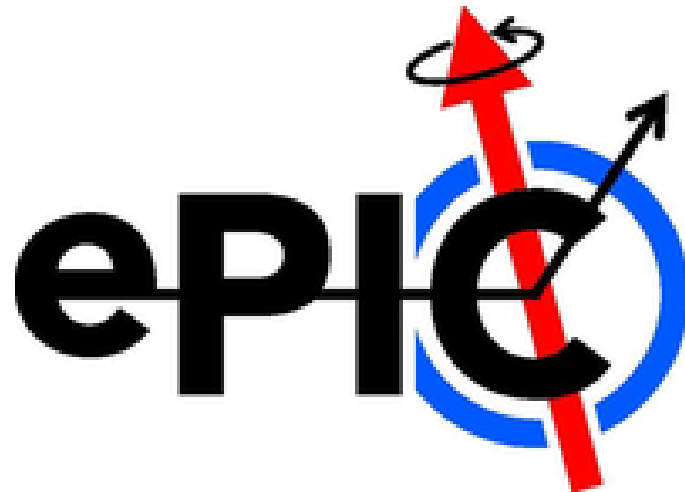


ePIC Italia

General Meeting

Update on data compression and SRO WG activities

Fabio Rossi



SRO WG Overview

ePIC DAQ uses streaming architecture

DAQ

Definition of streaming is “No L0 trigger”

- All data is zero suppressed by the front end electronics
- No system wide deadtime in normal operation
- Collaboration should have the full ability to make data selection cuts on the widest possible criteria
 - Flexible event selection, data selection and background characterization
- But subject to an overall throughput budget of ~100Gb/sec

ePIC Streaming will include

- Capabilities for hardware and firmware based triggering
- Capability for flow control
- Zero suppression & aggregation within data packets

Greater sensitivity to noise

Greater sensitivity to backgrounds

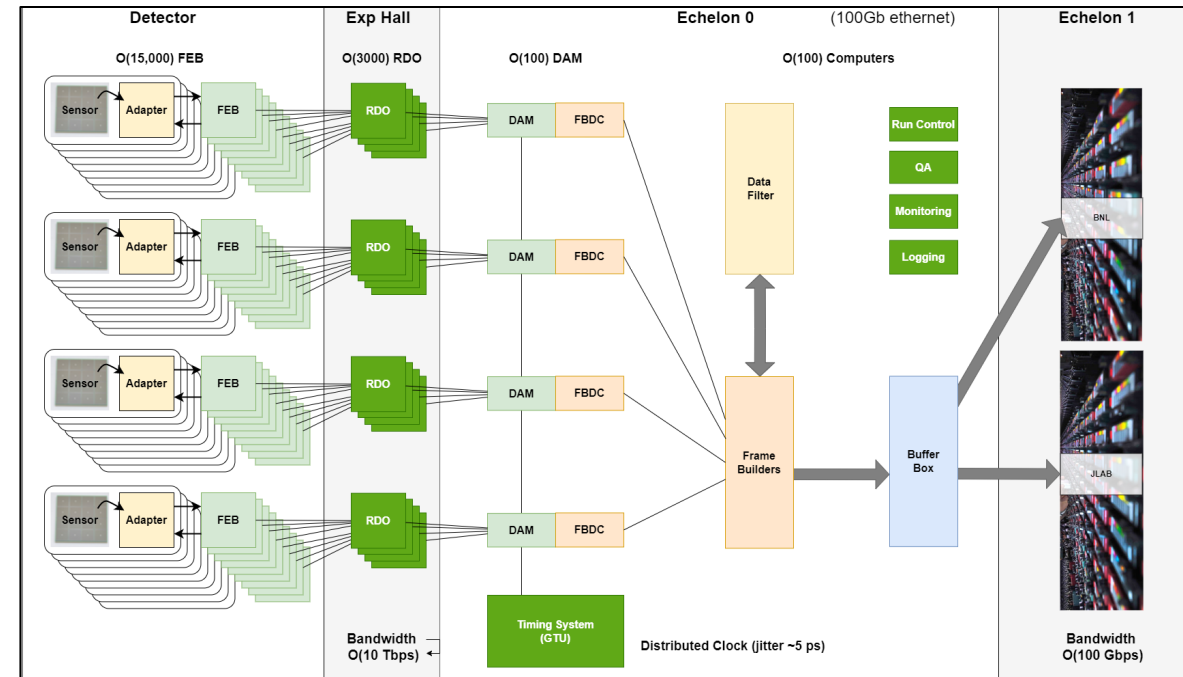
Computing

Definition of streaming is “Process data as it arrives”

- Fast Analysis (~3 weeks not months or years) using automation of calibration and reconstruction.

Requires some overlap between DAQ/computing

- Automation of calibrations
 - Mapping calibration dependencies
 - Mapping application of calibrations
 - Mapping evolution of calibrations
- QA and monitoring can make use of full offline structure
- Consistent schemes and language for data/metadata
- Event selection / tagging / and accounting

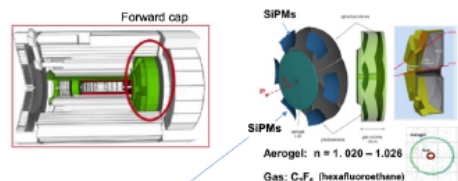


- Organizational contribution
- SRO Prototypes

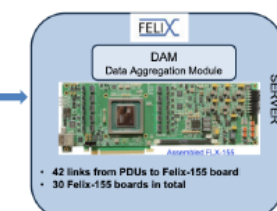
- Data reduction
- dRICH DAQ

dRICH

Dual Radiator RICH (dRICH)



PDU



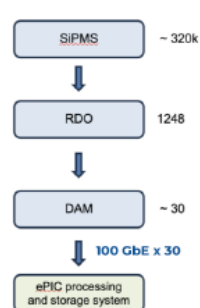
- 1 photodetector unit PDU: 4x64 SIPM array device (256 channels), 4 FEBs, 1 RDO
- 1248 PDUs for full dRICH readout
- 319488 readout channels divided in six sectors

- 42 links from PDUs to Felix-155 board
- 30 Felix-155 boards in total

4

Analysis of dRICH Output Bandwidth

The dRICH DAQ chain in ePIC → the throughput issue



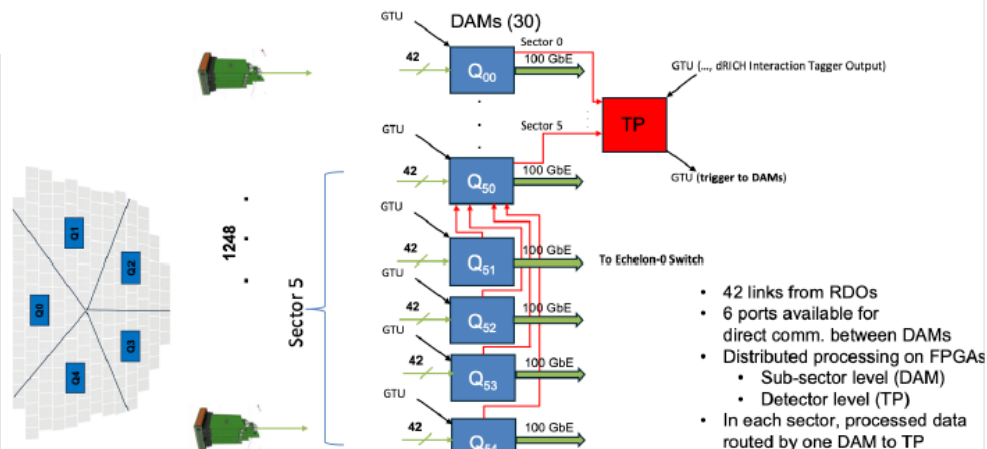
dRICH DAQ parameters		
RDO boards		1248
ALCOBA + RDO		4
dRICH channels (bits)		919488
Number of DAM L1		27
Input link in DAM L1		47
Output link in DAM L1		1
Number of DAM L2		1
Input link in DAM L2		27
Link bandwidth (1 Gbps) (assuming VTRON)		10
Interaction trigger reduction factor		1
Interaction trigger latency (s)		2.00E-04
ePIC parameters		
ePIC chips (chips)		66,522
Chip efficiency (taken into account gap)		0.90

Bandwidth analysis		
		Limit
Sensor rate per channel (bits)	300,000	4,000,000
Rate post-chiller (bits)	55,000	800,000
Throughput to FELIX (Mbps)	34,500	768,120
Throughput from ALCOBA (Mbps)	275,000	
Throughput from RDO (Mbps)	1,100	10,000
Input or each DAM (Mbps)	55,000	470,000
Buffering capacity at DAM (bits)	12,071	
Output from every DAM	7,400	10,000
Total throughput	1,344,240	270,000

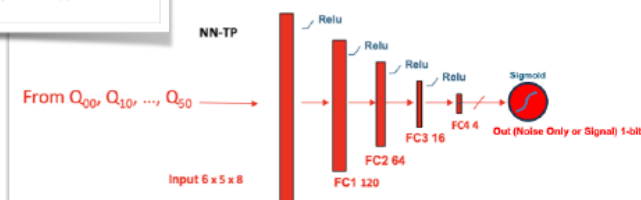
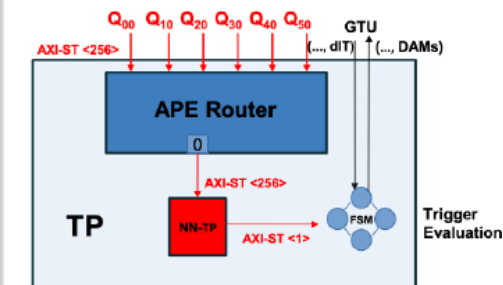
- Sensors DCR: 3-300 kHz (increasing with radiation damage → with experiment lifetime).
- Full detector throughput (FE): 14-1400 Gbps
- A reduction is needed to match 30 channels aggregated bandwidth (and safety margin)
- EIC beams bunch spacing: 10 ns → bunch crossing rate of 100 MHz
- For the low interaction cross-section (DIS) → one interaction every ~100 bunches
- A system tagging the (DIS) interacting bunches could solve the issue reducing down to ~1/100 the data throughput

- Two complementary approaches are possible:
1. Develop a dedicated sub-detector tagging relevant interactions.
 2. This proposal.

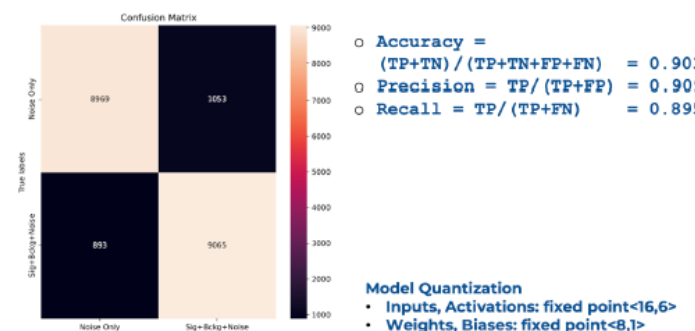
dRICH DAQ & Data Reduction



- 42 links from RDOs
- 6 ports available for direct comm. between DAMs
- Distributed processing on FPGAs
 - Sub-sector level (DAM)
 - Detector level (TP)
- In each sector, processed data routed by one DAM to TP

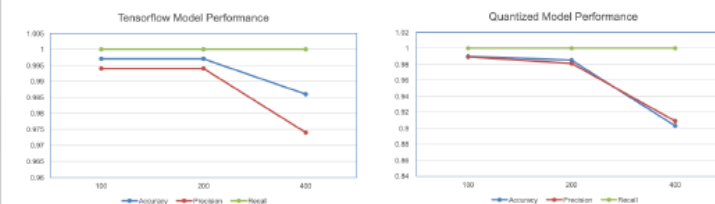


Quant. model performance @ noiseRate = 400 KHz



- Accuracy = $\frac{TP+TN}{TP+TN+FP+FN} = 0.903$
- Precision = $\frac{TP}{TP+FP} = 0.909$
- Recall = $\frac{TP}{TP+FN} = 0.895$

Summary of Distributed MLP Performance



Realistic Noise Model

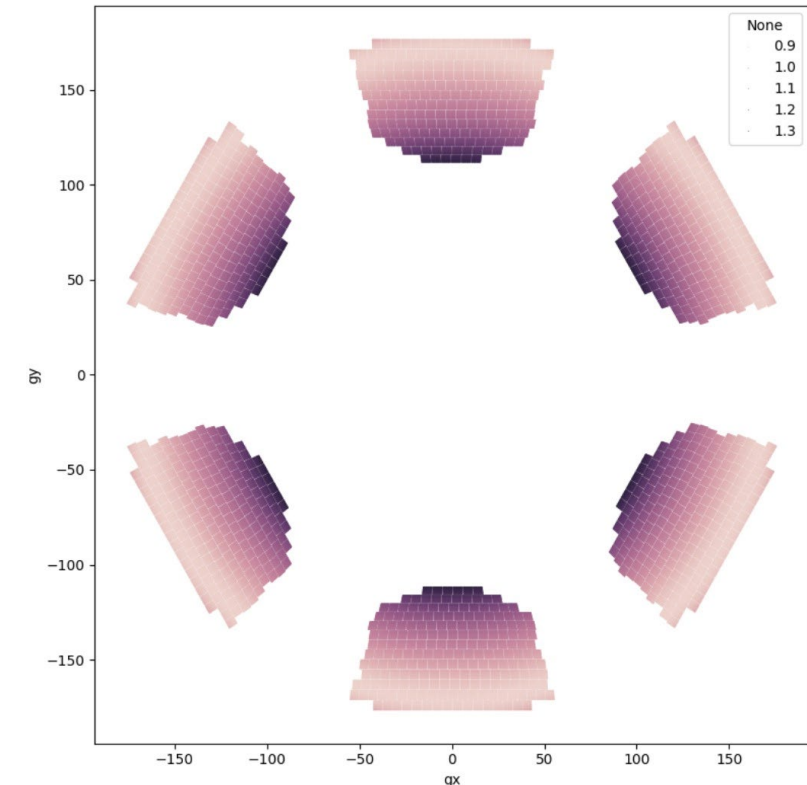
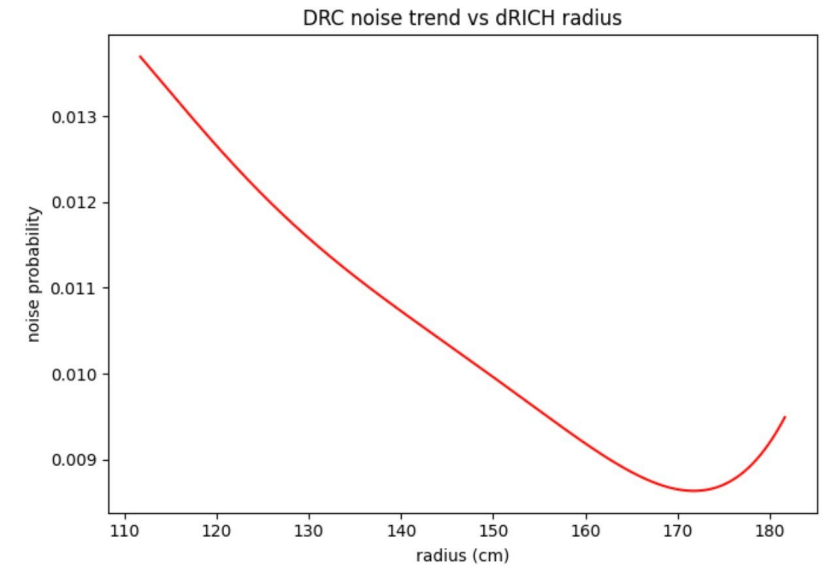
Thanks to R. Preghenella's contribution, we added to the toy noise model ("gaussian") a new "realistic" one in which the Dark Count probability of a certain dRICH SiPM is dependent on its radial distance from the detector z-axis and on the integrated luminosity

⇒ Implemented in EICRecon digitization step
(new flag to enable new model noise)

```
const float baseline_dcr = 3.e3; // [Hz] new sensors at T = -30 C and Vover = 4V
const float dcr_increase = 300.e3 / 1.e9; // [Hz/neq]
float neq_radius_params[6] = { -3.27029e+09, 1.26055e+08, -1.88568e+06, 13929.1, -50.9931, 0.0741068 };

float neq_radius(float radius /* cm */)
{
    float neq = 0.;
    for (int ipar = 0; ipar < 6; ++ipar)
        neq += neq_radius_params[ipar] * std::pow(radius, ipar);
    return neq;
}

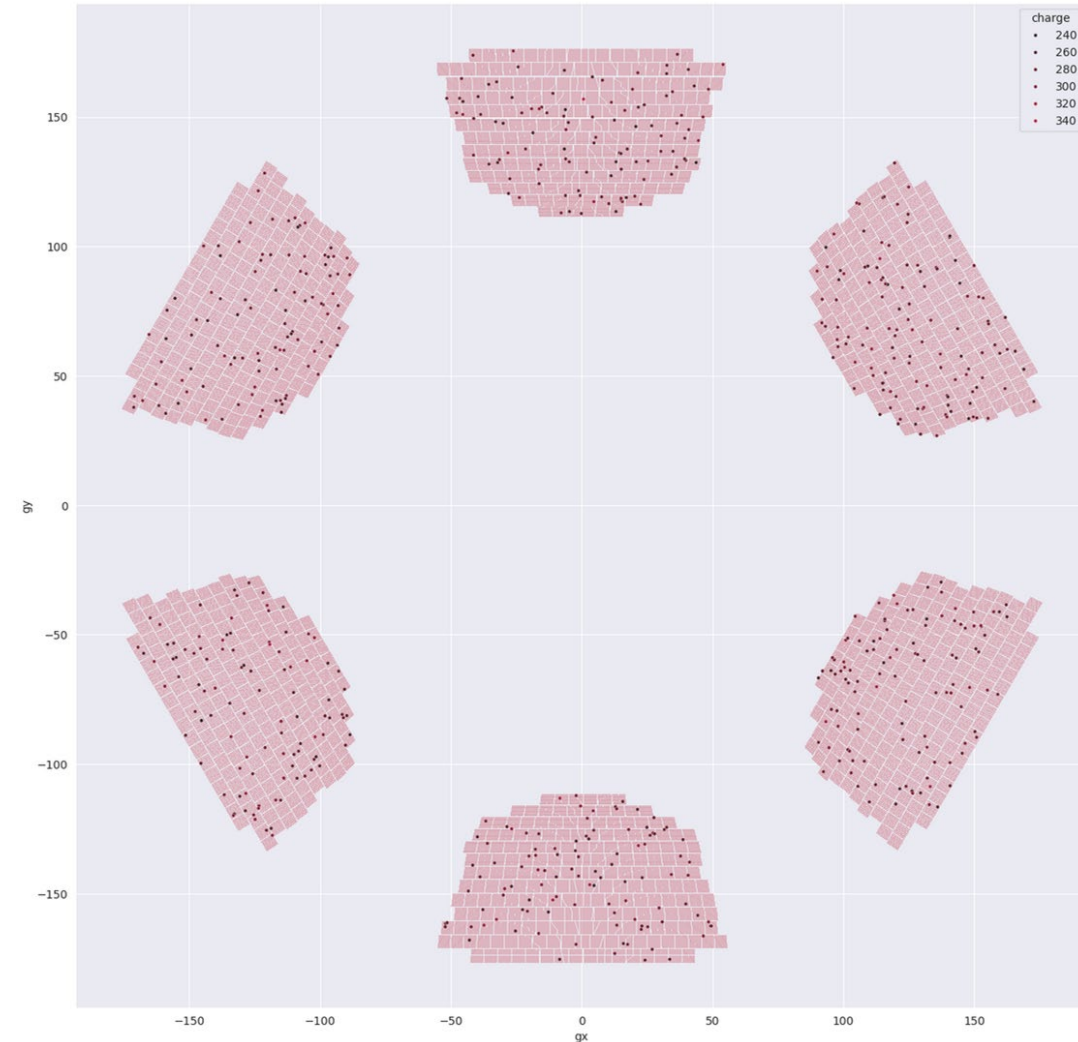
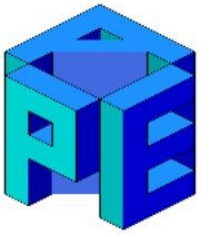
float
noise_probability(float radius = 150. /* cm */, float window = 10. /* ns */, float luminosity = 100. /* fb-1 */)
{
    float neq = neq_radius(radius) * luminosity;
    float dcr = baseline_dcr + dcr_increase * neq;
    float pro = dcr * window; /* 1.e-9;
    return pro;
}
```



Reconfigurable parameters, used for the new dataset:

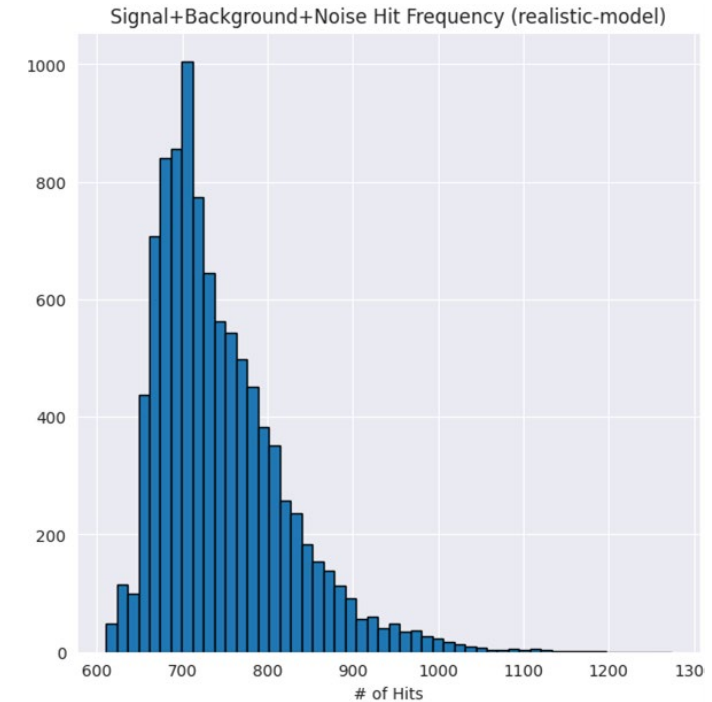
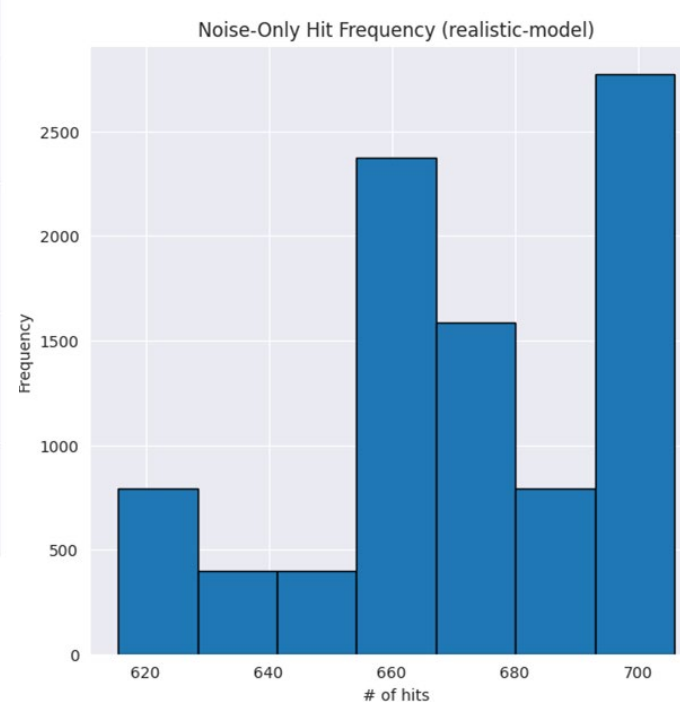
- (Noise) Time window = 10 ns
- (Integrated) Luminosity = 2 fb-1

Updates to the Data Reduction Pipeline

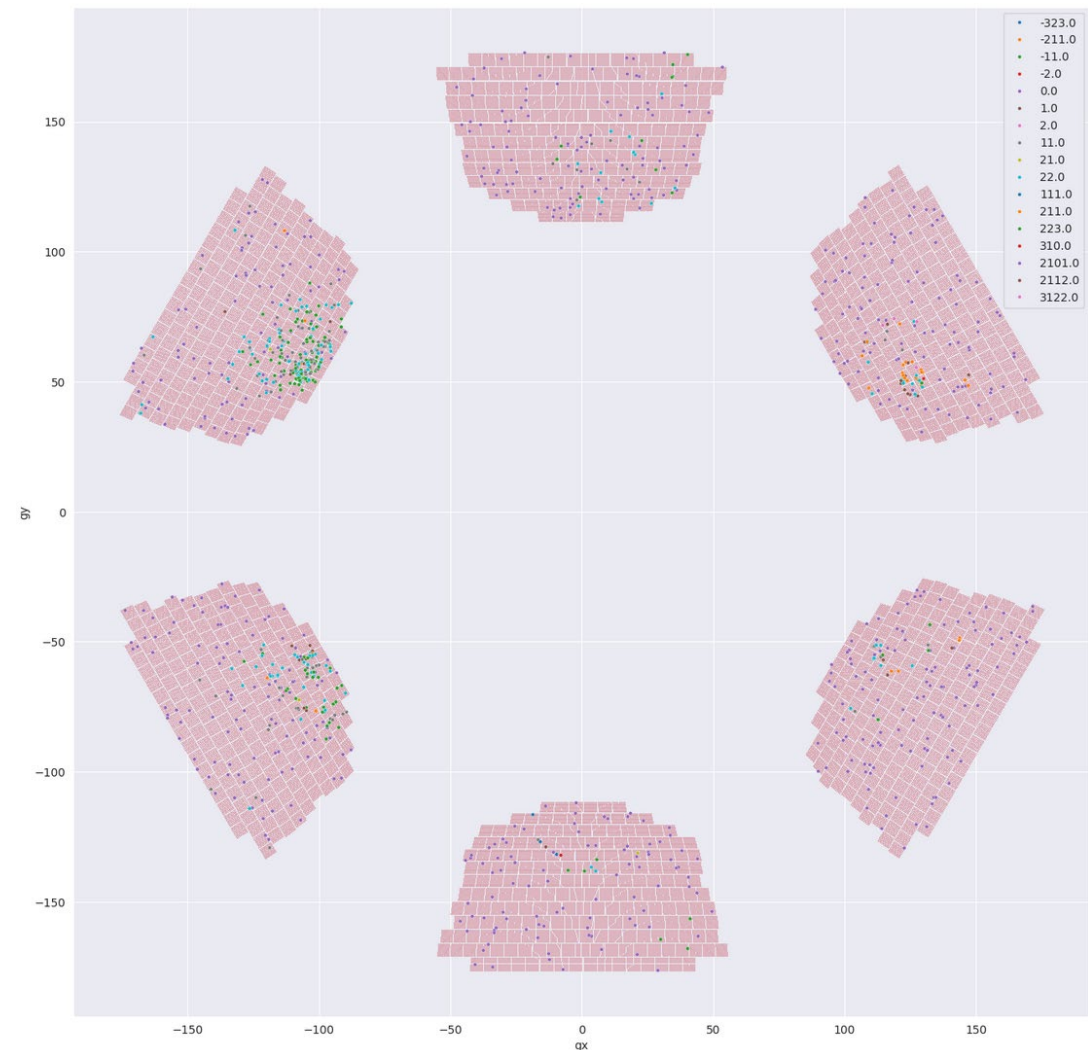
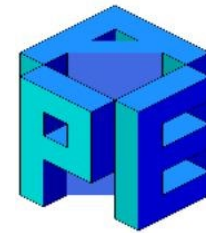


With the new “realistic” noise model, the **global density of dark count noise hits** within the selected events has increased with respect to the previous dataset used for training and testing.

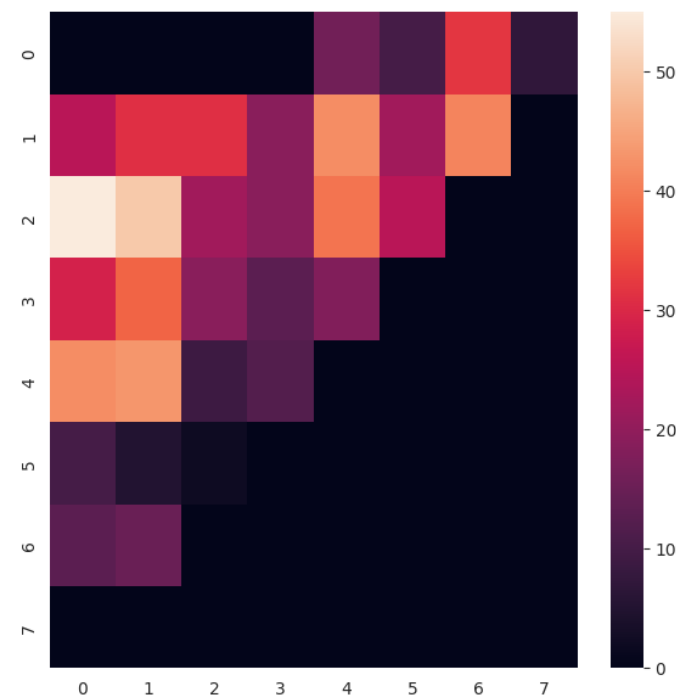
⇒ we decided to maintain the same NN model architecture (in order to cope with the hardware constraints) and to re-train the whole model with the new dataset



Updates to the Data Reduction Pipeline

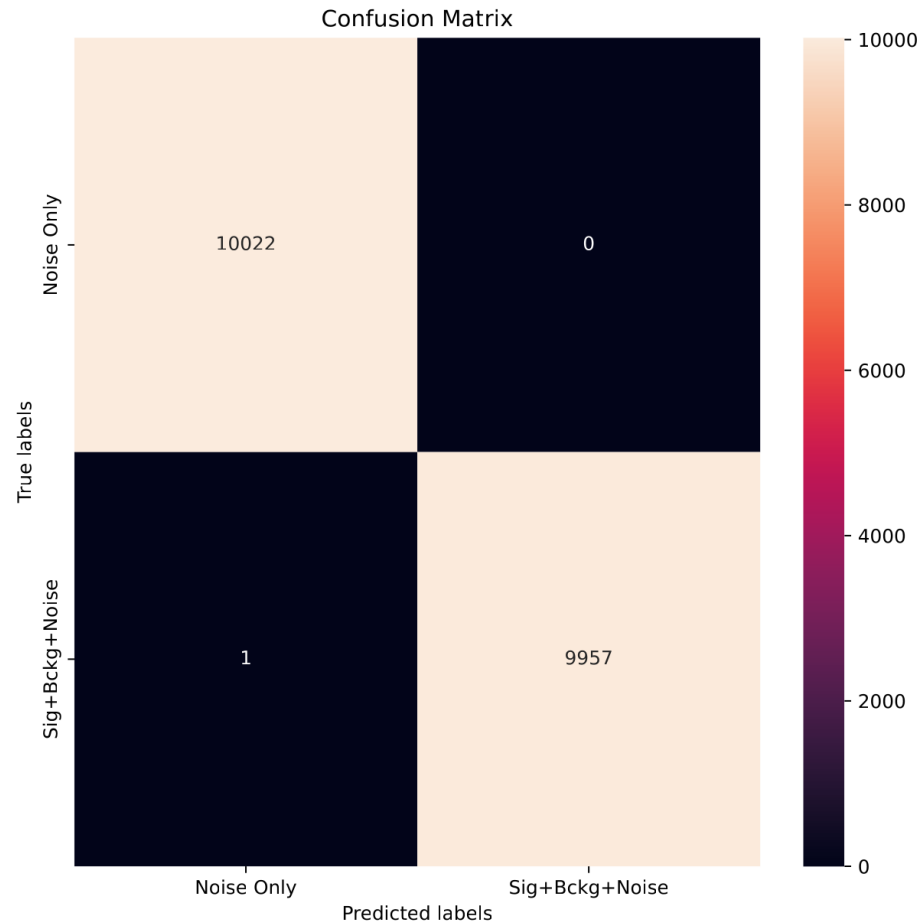
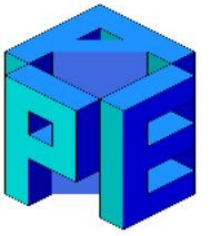


⇒ taking as example a **background event**, we can see how the **pixel intensity** in input to the **MLP_DAM NN** has increased wrt to the previous cases



⇒ we explore several new preprocessing steps (rescaling, normalization,...) to take care about this new feature in the data

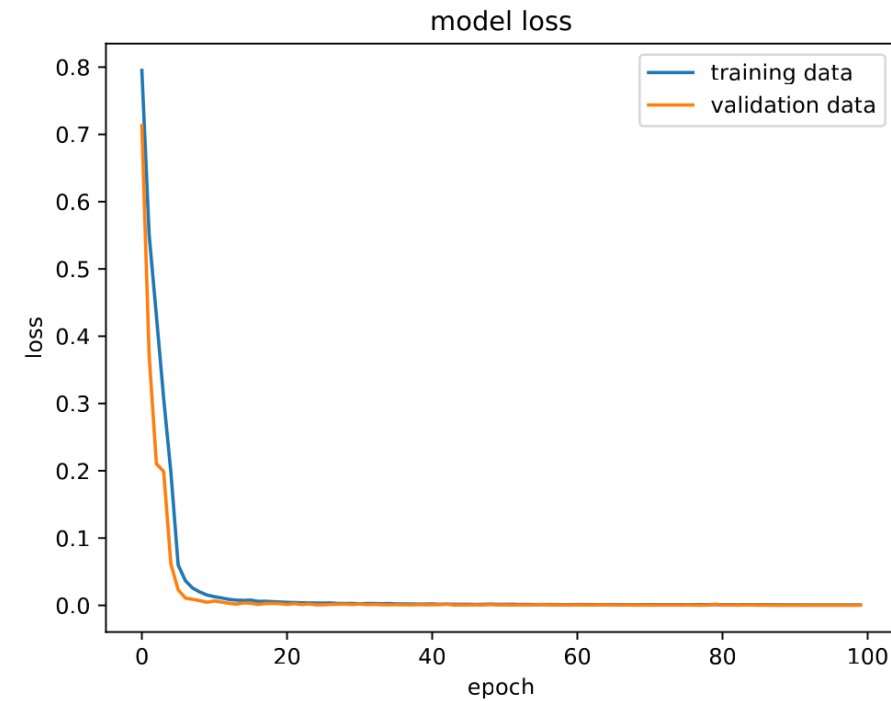
Results with the new noise model & pipeline



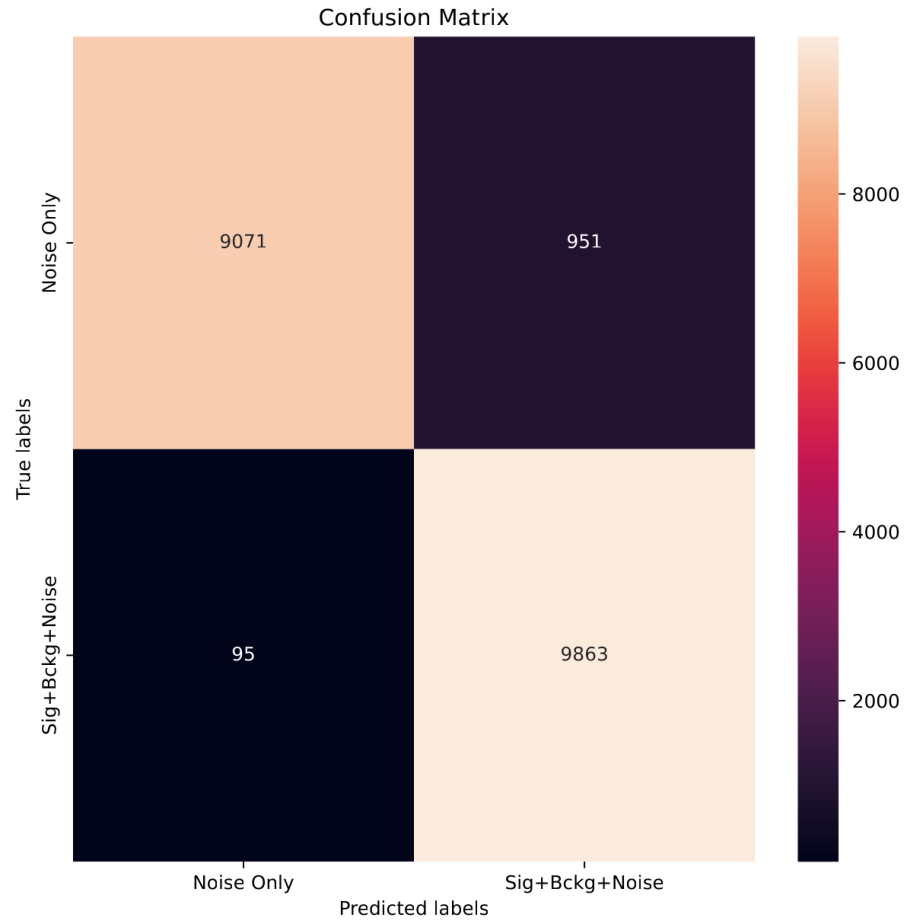
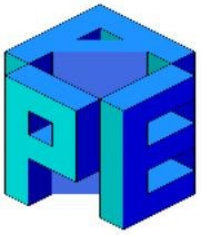
Accuracy = $(TP + TN) / (TP + TN + FP + FN) = 0.9999$

Purity = $TP / (TP + FP) = 0.9999$

Efficiency = $TP / (TP + FN) = 1.0000$



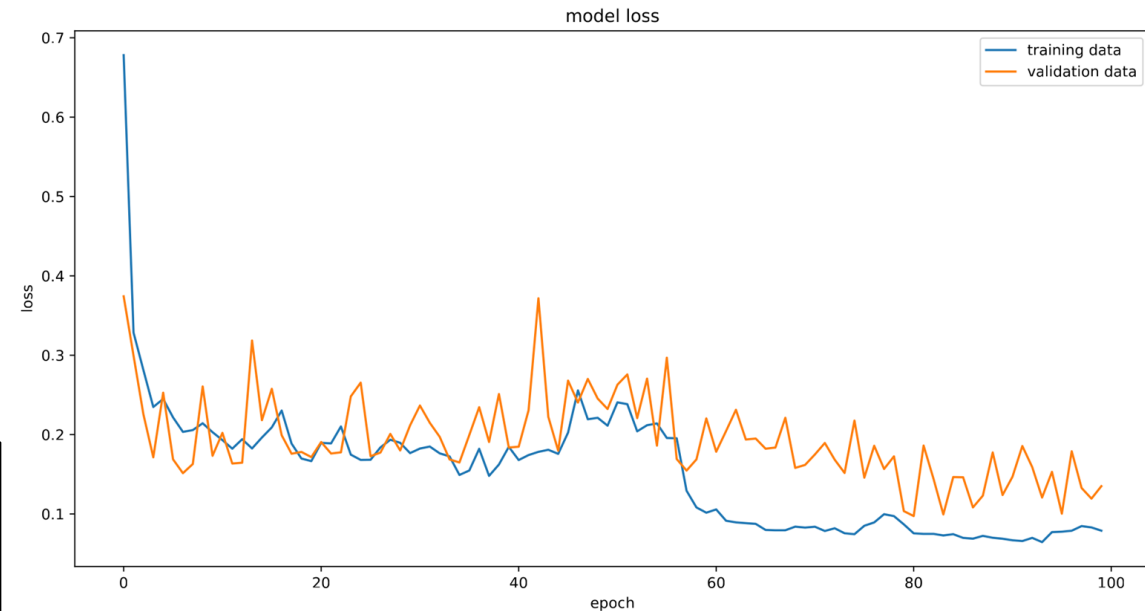
Results with the new noise model & pipeline



Accuracy = $(TP + TN) / (TP + TN + FP + FN) = 0.9051$

Purity = $TP / (TP + FP) = 0.9896$

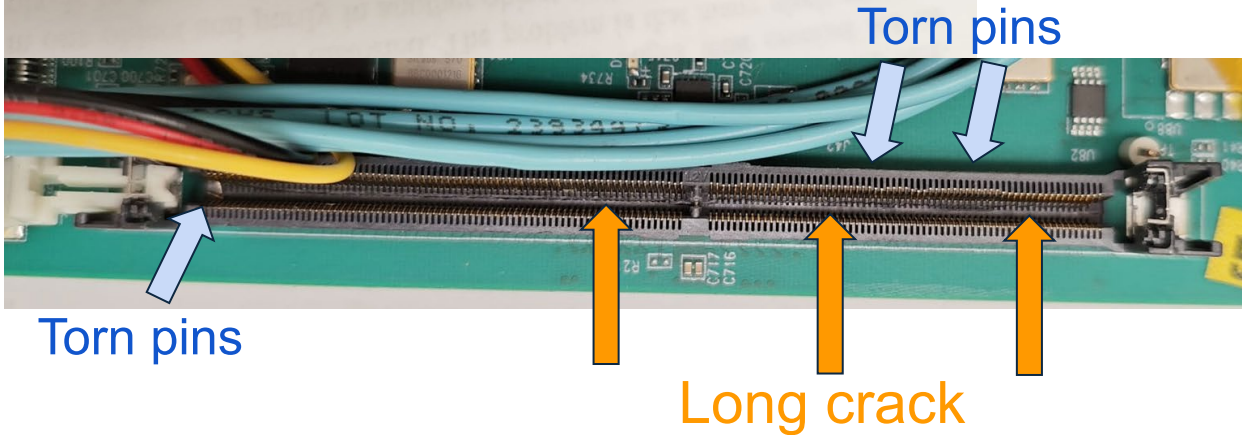
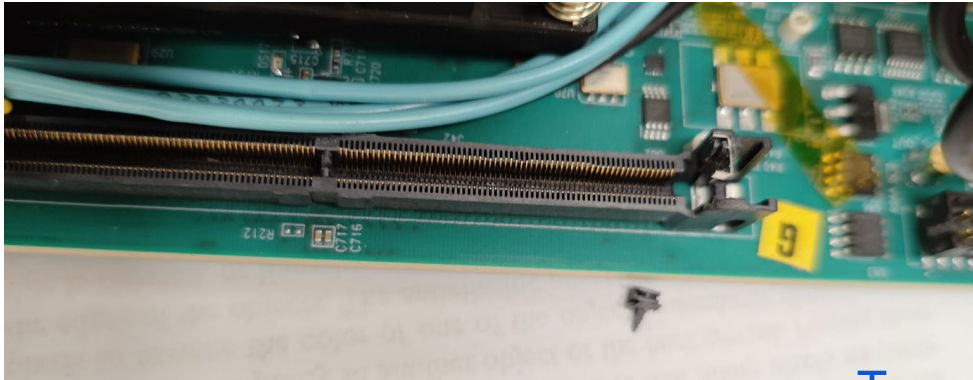
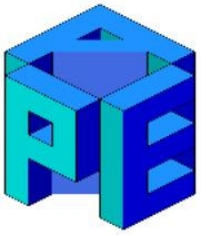
Efficiency = $TP / (TP + FN) = 0.9051$



Model Quantization:

- Inputs, Activations: fixed point<16,6>
- Weights, Biases: fixed point<8,1>

Procurement of one Felix-182 Card



Felix-182 board arrived from JLab in Rome end of December '24
Unfortunately it showed damages to the DRAM slot:

- **torn contact pins**
- **a crack along the inner side of the slot toward the FPGA**

Maybe they occurred due to the pressure of a DRAM module left in the slot during the delivery (bad packaging).

As a consequence, DRAM is not detected by the system.

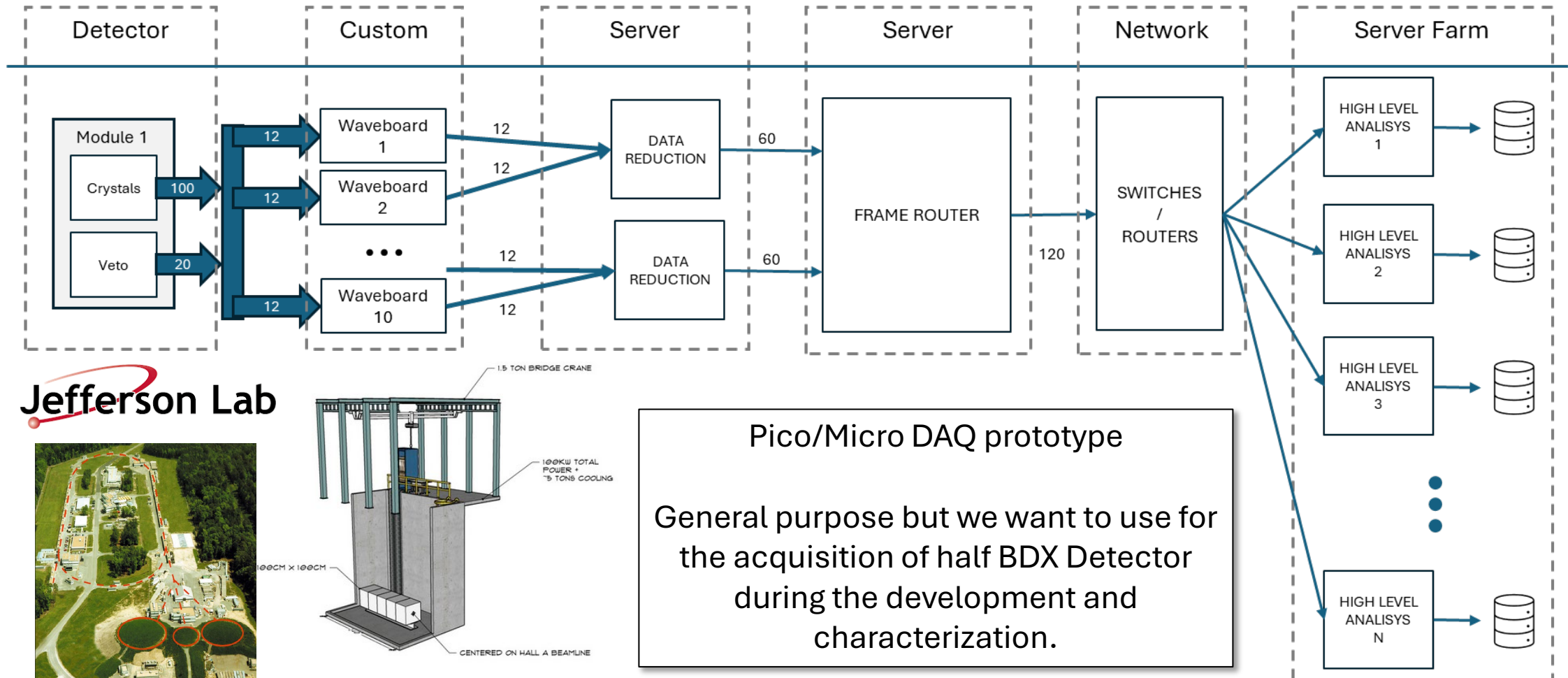
We tried to load some standard fw to the board but it was not possible because the DRAM check failed.

We consulted with:

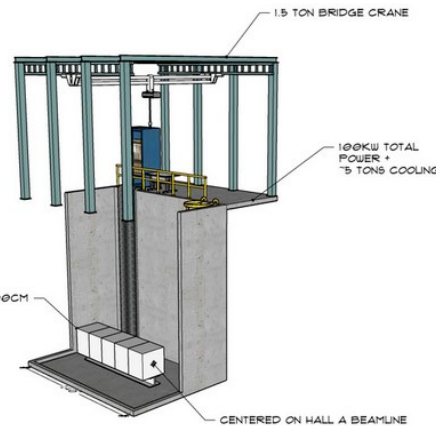
- the INFN Electronics Laboratory in Rome
- CERN EP-ESE Electronic Systems for Experiments (thanks to Markus Joos)

Both concur that is extremely difficult to repair the DRAM slot and not worth the effort considering the costs and the likely not optimal result, we are sending it back.

Activities in Genova



Jefferson Lab



Beam Dump eXperiment

Pico/Micro DAQ prototype

General purpose but we want to use for the acquisition of half BDX Detector during the development and characterization.

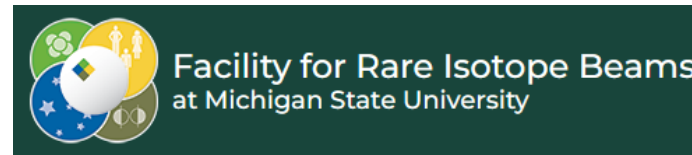
Organizational contribution

SPADI Alliance

Signal processing and data acquisition infrastructure alliance

Jefferson Lab

Brookhaven
National Laboratory



We all share the Same Goal
Develop a general-purpose streaming
acquisition system.

Streaming Data Acquisition System Meeting

MSU - JLab

Streaming Readout Workshop

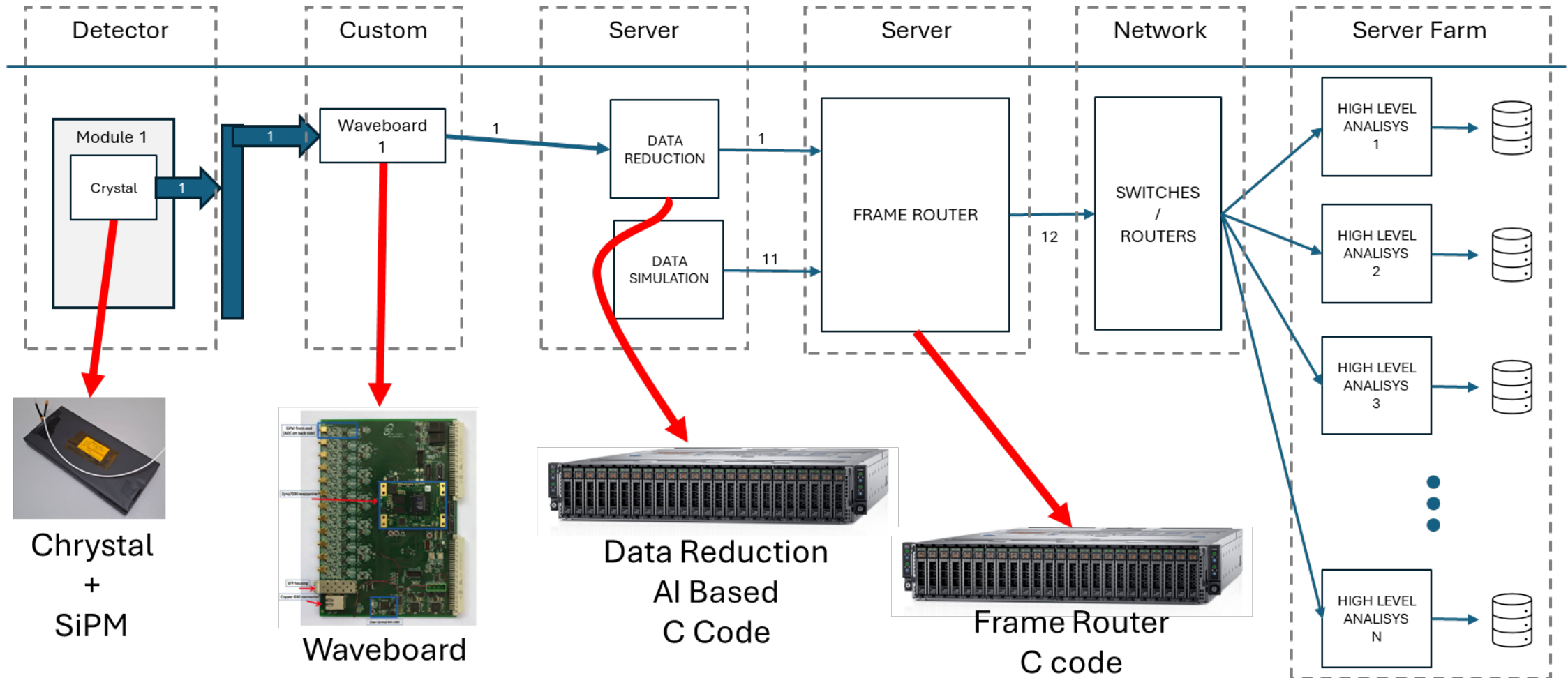
SRO-XII

December 2-4, 2024
University of Tokyo, Japan



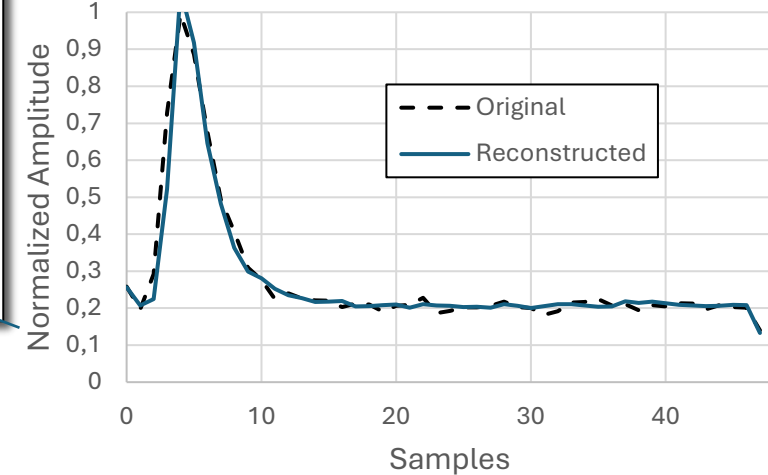
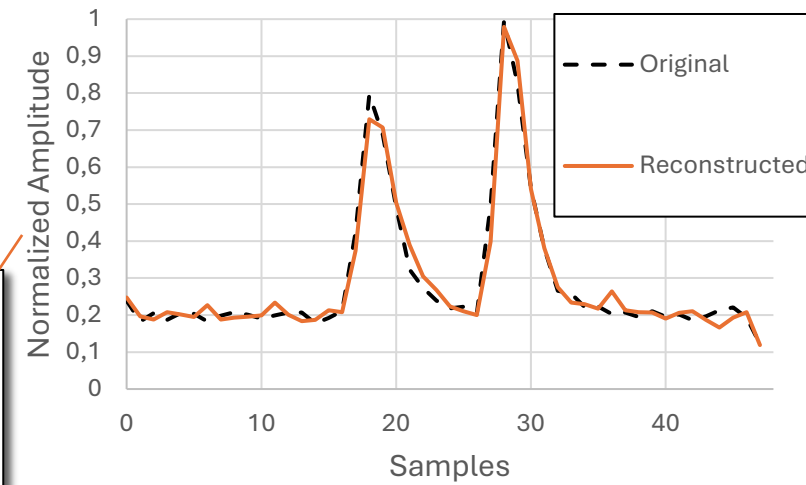
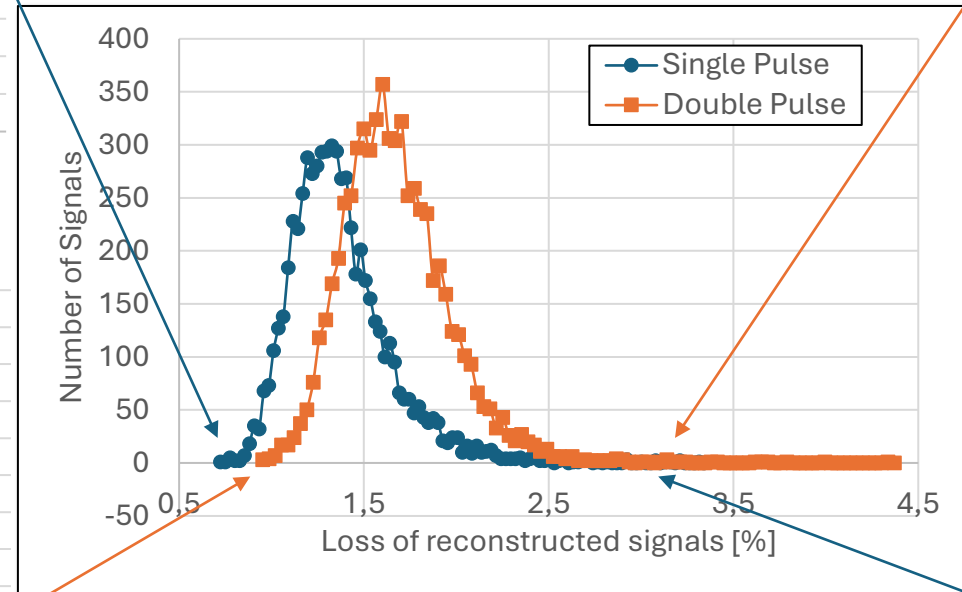
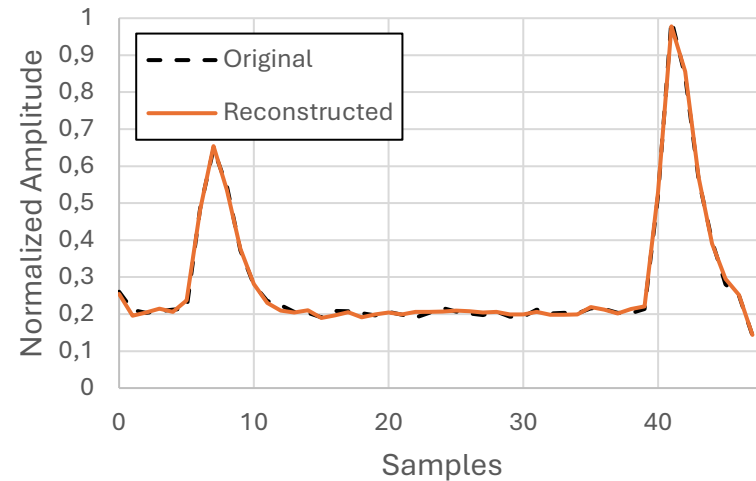
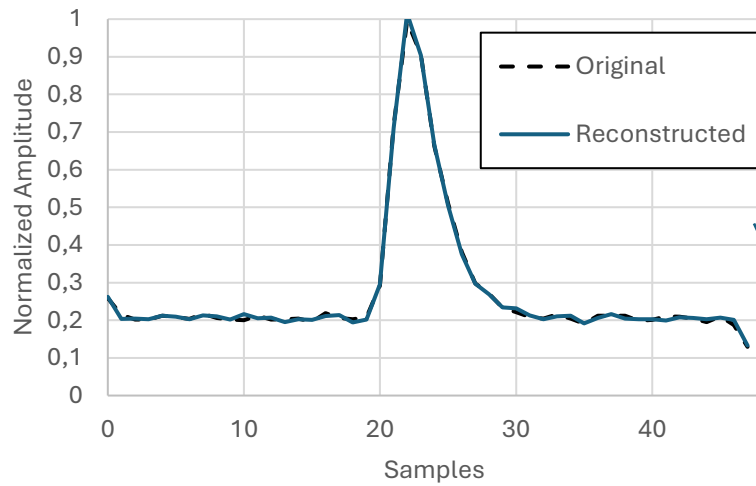
February 12, 2025
Jefferson Lab, Newport News (VA)

BDX test Setup: state of work



Data Reduction

Test results for CLAS12 signals



Lossy compression based on AI algorithm
Compression ratio: 4
Introduced loss: $\approx 2\%$



Finanziato
dall'Unione europea
NextGenerationEU



Ministero
dell'Università
e della Ricerca



Italiadomani
PIANO NAZIONALE
DI RIPRESA E RESILIENZA



Future
Artificial
Intelligence
Research

Data reduction: Implementation

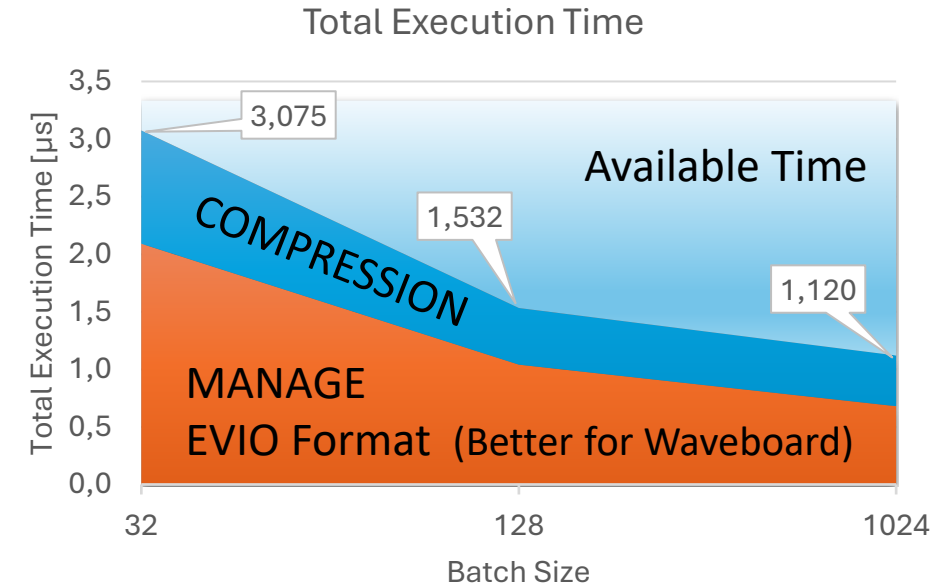
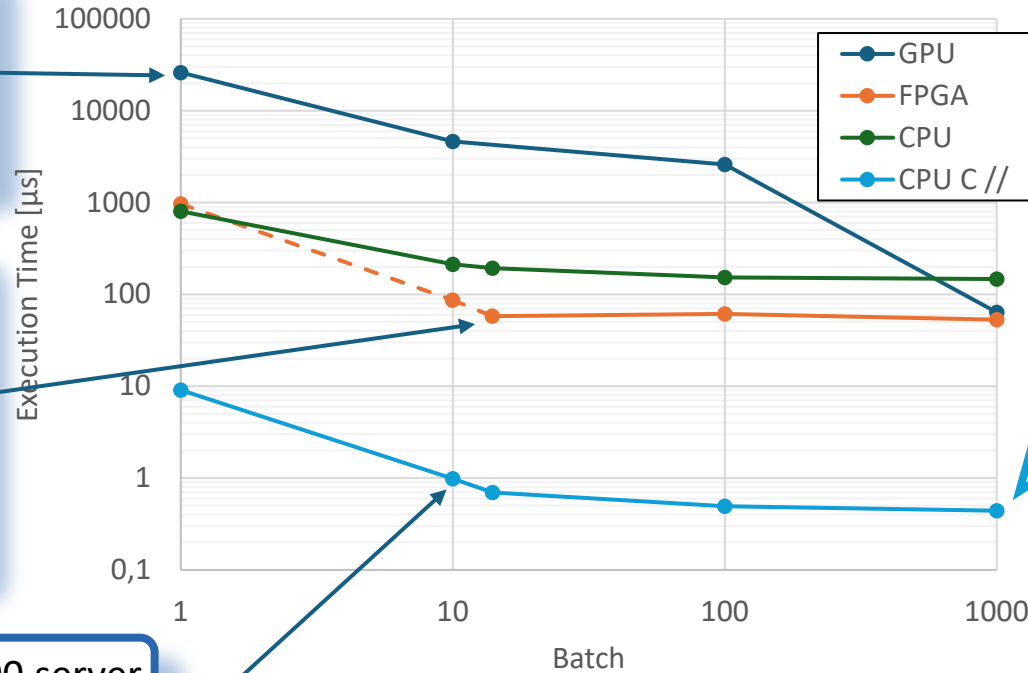
NVIDIA Tesla V100 GPU



Xilinx
XRT



ALVEO V70 FPGA



Test results for CLAS12 signals

AMD presentation @ Workshop on Electronics Torino:
<https://agenda.infn.it/event/44098/contributions/253506/>

Next Step

Hardware level

HLS4ML

Dedicated connectivity



Alinx VD 100



Finanziato
dall'Unione europea
NextGenerationEU



Ministero
dell'Università
e della Ricerca

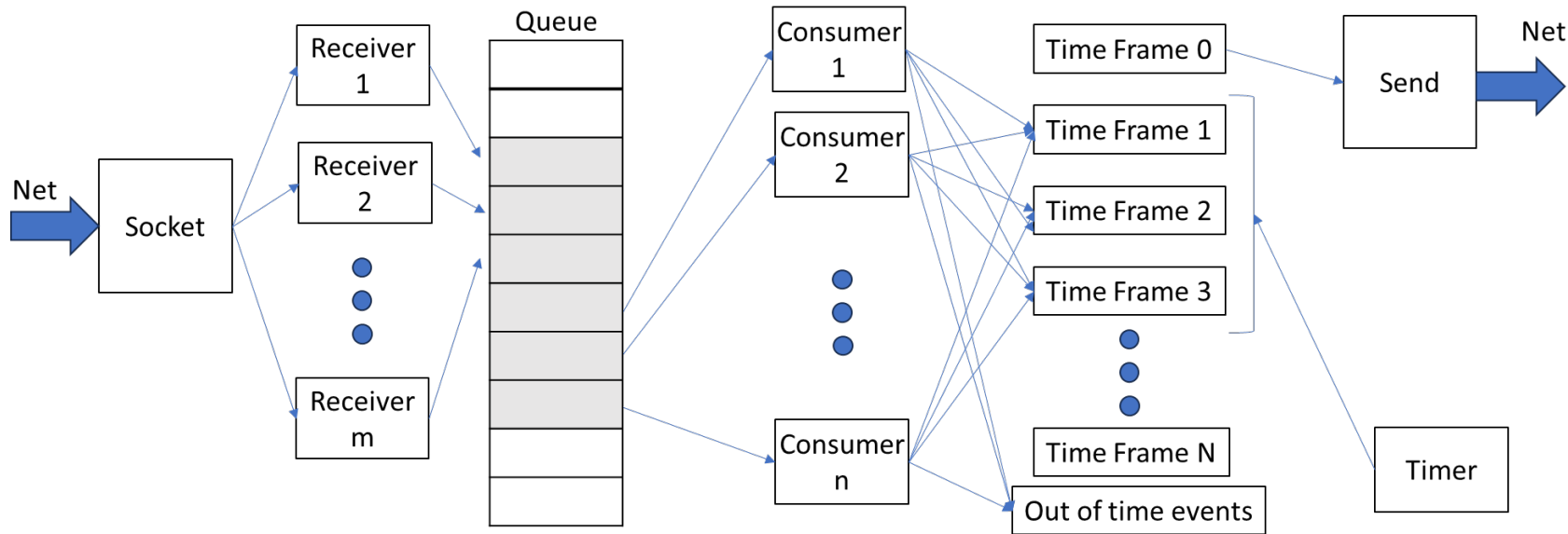


Italiadomani
PIANO NAZIONALE
DI RIPRESA E RESILIENZA



Future
Artificial
Intelligence
Research

Frame Router



Details:

- C code
- multithread by means of OpenMP
- CPU Server implementation

First Version are almost
done for testing



Required Steps:

Get Full control of Waveboard

Use Multiple Waveboard

Write Custom Run Control

Test and performance evaluation

Thank you for your attention