

Intelligent Experiments Through Real-Time AI

*Fast Data Processing and Autonomous Detector Control
for sPHENIX and Future EIC Detectors*

Ming Liu
Los Alamos National Laboratory
for the Fast-ML Team

DOE Fast-ML Presentations
November 30, 2022

Today's Presentations

1. Overview, 10'

- Ming Liu (LANL)/Gunther Roland(MIT)/Nhan Tran(FNAL)/ Dantong Yu (NJIT)

2. Physics simulation and AI-ML algorithms, 10'

- Dantong Yu (NJIT)/Cameron Dean(MIT)/Zhaozhong Shi(LANL) /Hang Qi(MIT)/Hao-Ren Jheng(MIT)/Beilei Jiang(NTU)

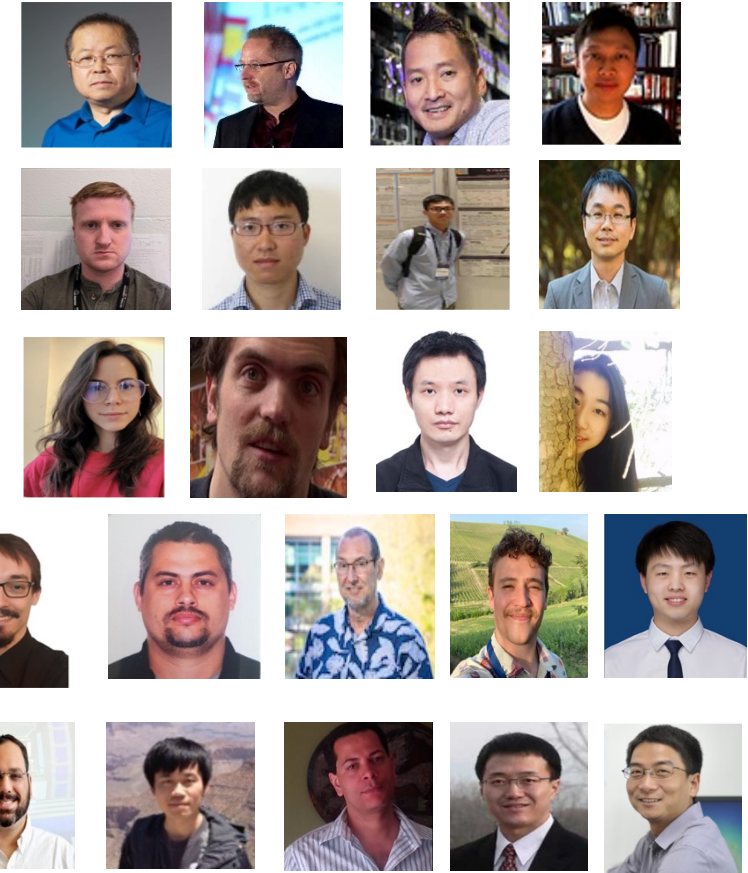
3. HLS4ML and firmware implementation, 5'

- Micol Rigatti (FNAL)/Phil Harris(MIT)/Nhan Tran(FNAL)/

4. Demonstrator implementation, 5'

- Jakub Kvapil (LANL)/Yasser Corrales(MIT)/Noah Wuerfel(LANL)/Jo Schambach(ORNL)/Kai Chen(CCNU)/Lang Lei(CCNU)/Beilei Jiang(NTU)

Q & A: 5'

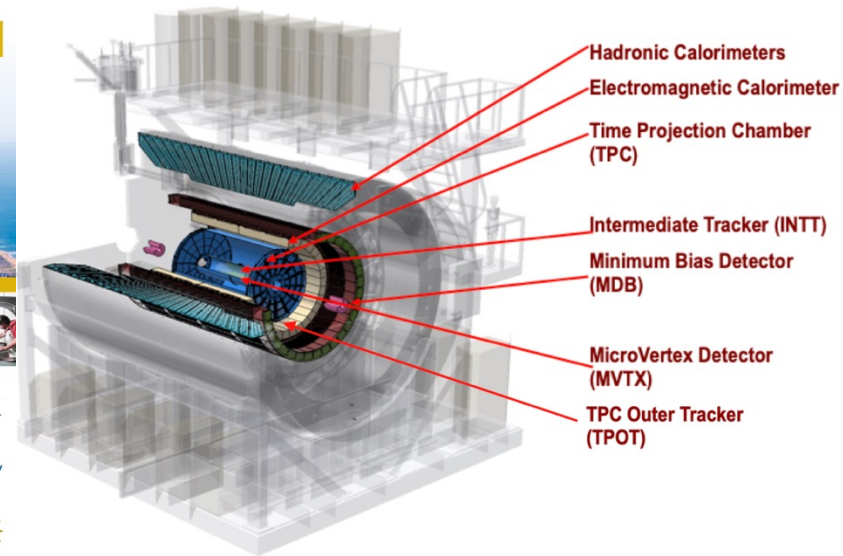




Overview

- Ming Liu

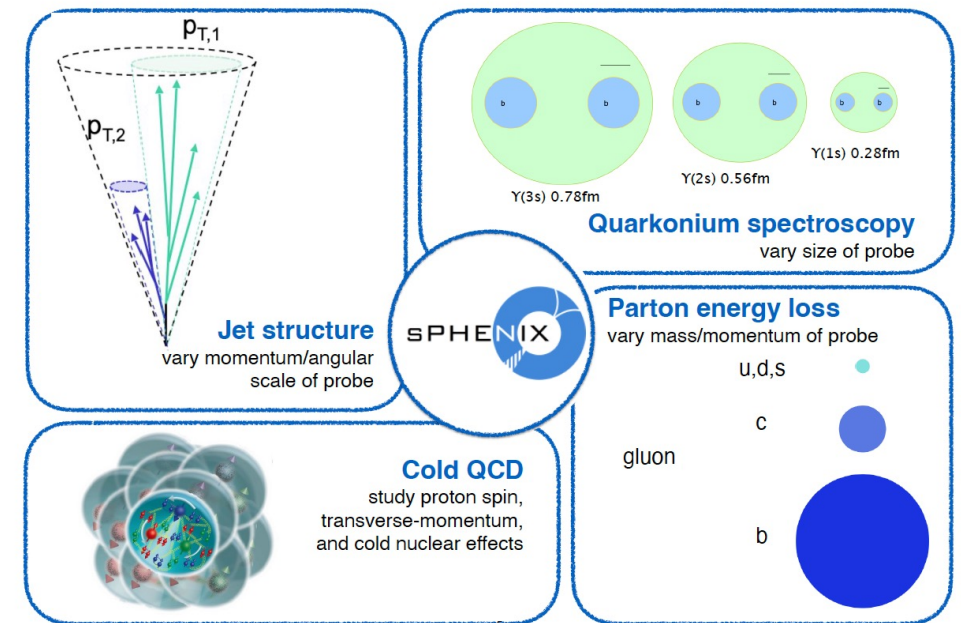
sPHENIX at RHIC



2015 NSAC Long Range Plan for Nuclear Science priority: sPHENIX Experiment at RHIC

- Probe the inner workings of QGP by resolving its properties at shorter and shorter length scales
- Complementary to LHC experiments to study relativistic heavy-ion collisions

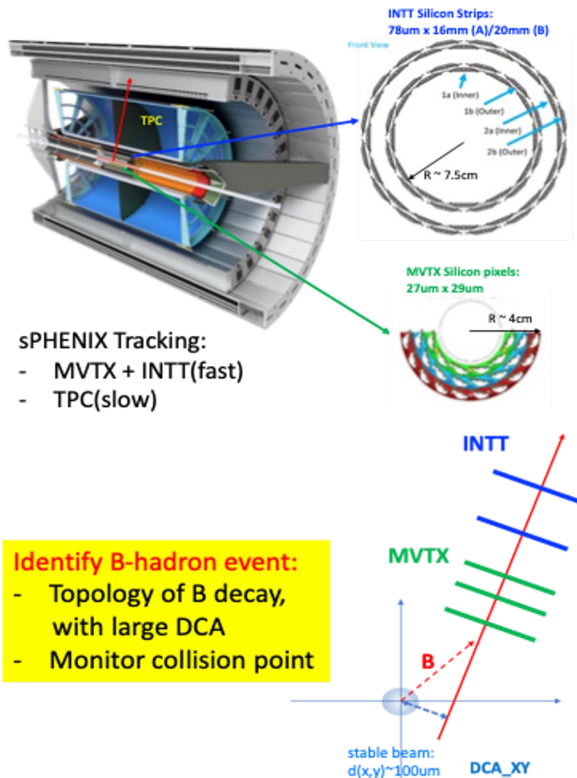
Heavy Flavor physics – a key pillar of RHIC science



Project Goals and Deliverables

Selective streaming real-time AI and autonomous detector control:

Deliver a demonstrator for p+p and p+Au running for sPHENIX -> generalizable for applications in experiments at the EIC, **improve b-hadron physics in p+p/Au by >100x**



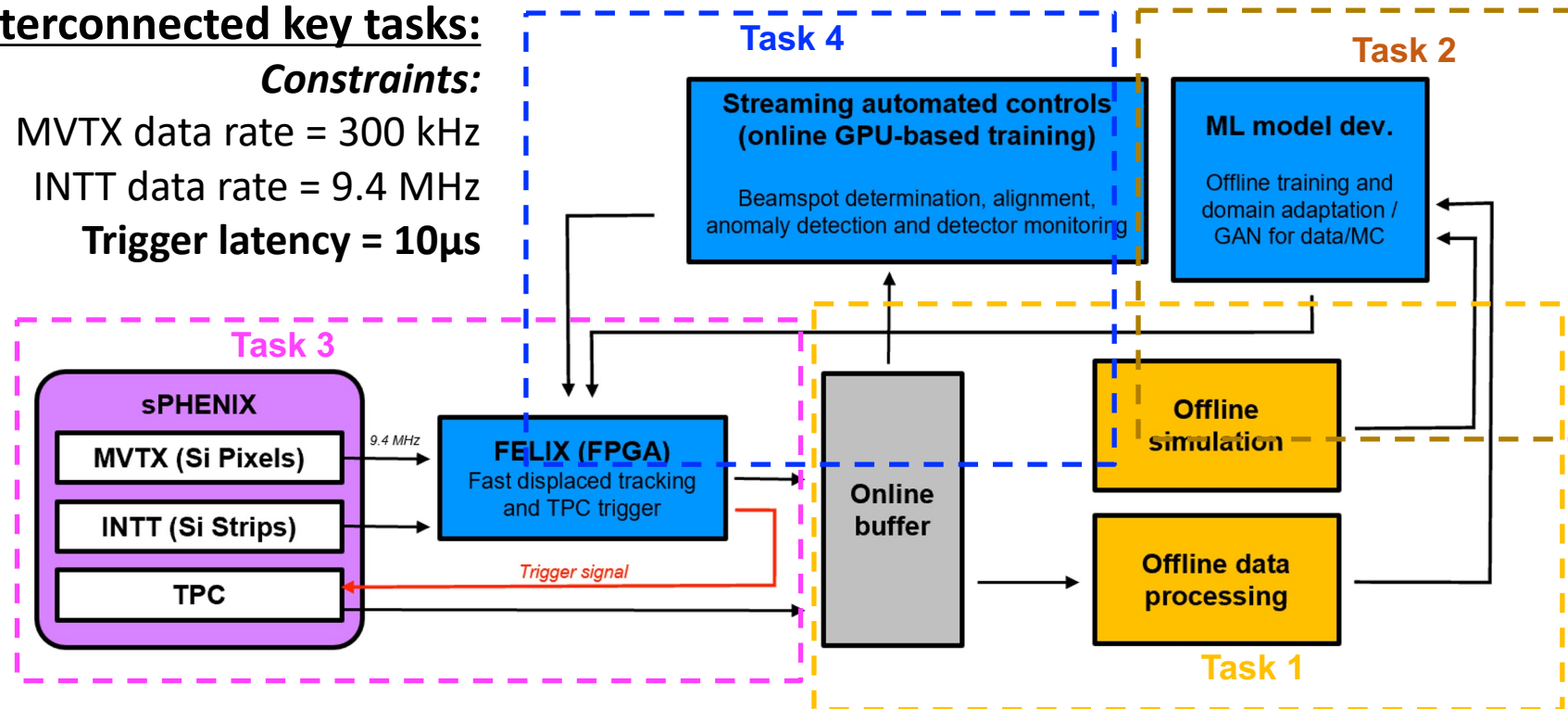
4 interconnected key tasks:

Constraints:

MVTX data rate = 300 kHz

INTT data rate = 9.4 MHz

Trigger latency = 10μs



Leadership and Technical Roles

Team of NP + HEP + CS

Los Alamos National Laboratory, Ming Xiong Liu, (**Lead Principal Investigator**)

Fermi National Laboratory, Nhan Tran, (**Co-PI**)

Massachusetts Institute of Technology, Gunther M Roland (**Co-PI**)

New Jersey Institute of Technology, Dantong Yu (**Co-PI**)

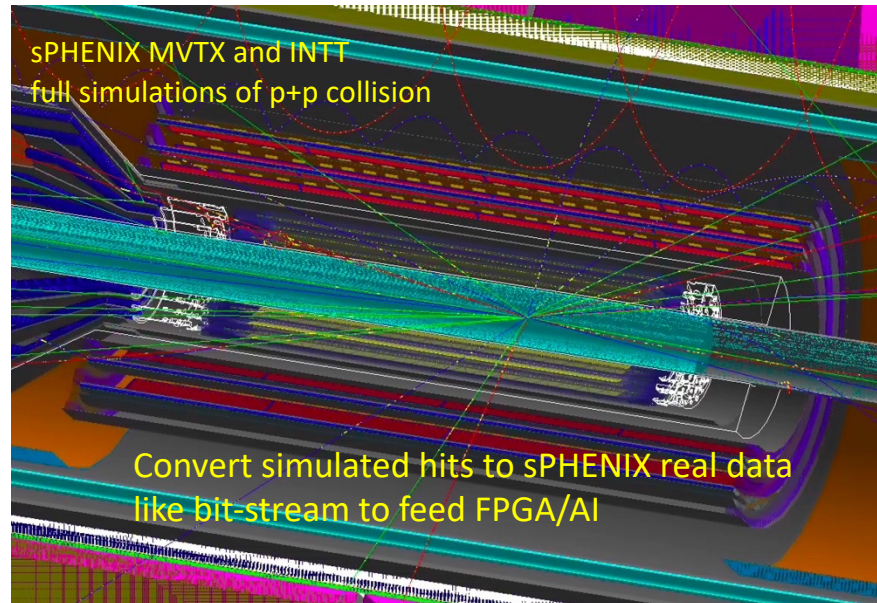
Leadership structure of the team The project team will be led by Lead Principal Investigator, Dr. Ming Xiong Liu of LANL, who is accountable to the DOE program leadership for the project's overall success. The team shares the responsibility and accountability for success. Within that structure, lead roles are assigned to co-Principal Investigators (co-PIs), also referred to as key personnel. Dr. Liu will be the lead for hardware design. Dr. Gunther Roland will be the physics lead in sPHENIX and EIC. Dr. Tran will be the lead for Co-Design of AI software and Hardware. Dr. Yu will be lead for Deep Neural Networks Software Design.

New teams joined later in early 2022:

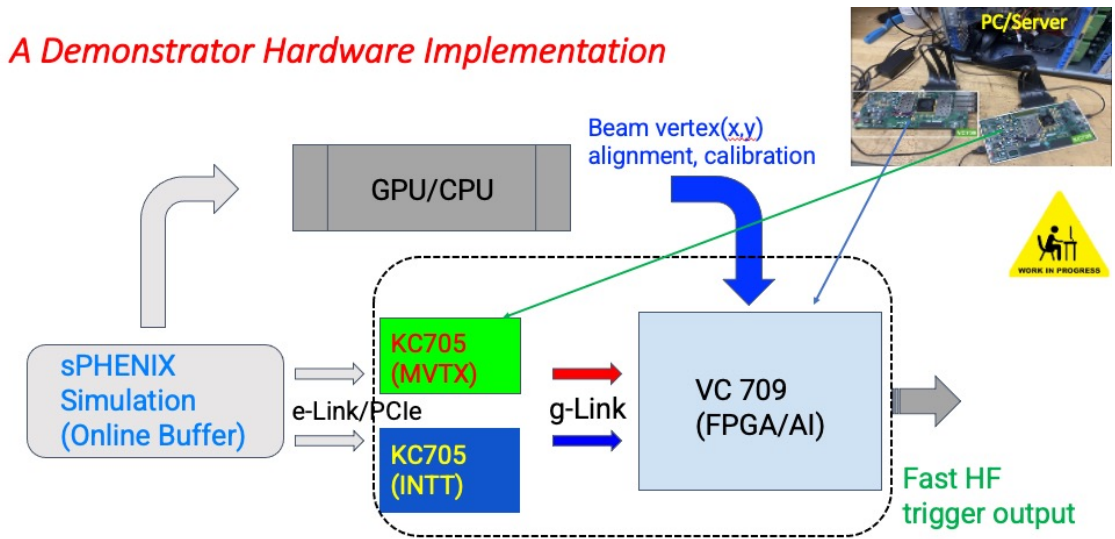
- Dr. Jo Schambach, ORNL, sPHENIX/EIC readout integration, sPHENIX MVTX and EIC/ePIC readout lead
- Dr. Kai Chen, CCNU, FELIX-AI-Trigger hardware integration, FELIX developer at BNL for ATLAS, also sPHENIX
- Prof. Song Fu, NTU, data acceleration in ML

Technical Approaches and Highlights - I

- **Objective 1 – Design, build, simulate, and benchmark a prototype streaming readout system with AI-based fast online data processing and autonomous detector control system that meets the physics and engineering requirements.** To support this objective, we first aim to generate a large volume of simulation data for heavy flavor decay events. We plan to design a prototype in the simulated and the real sPHENIX experimental environment and later apply the technology in the high luminosity EIC experiments at RHIC. Our objective is to create a working prototype that serves as a baseline and template for future upgrades. With this prototypical working solution, we target to improve the heavy flavor samples from the current 0.05% yield to more than 10+%. **(Task 1)**



A Demonstrator Hardware Implementation



Streaming readout simulation data:

8b/10b MVTX/INTT data (KC705) to FPGA/AI Engine (VC709)

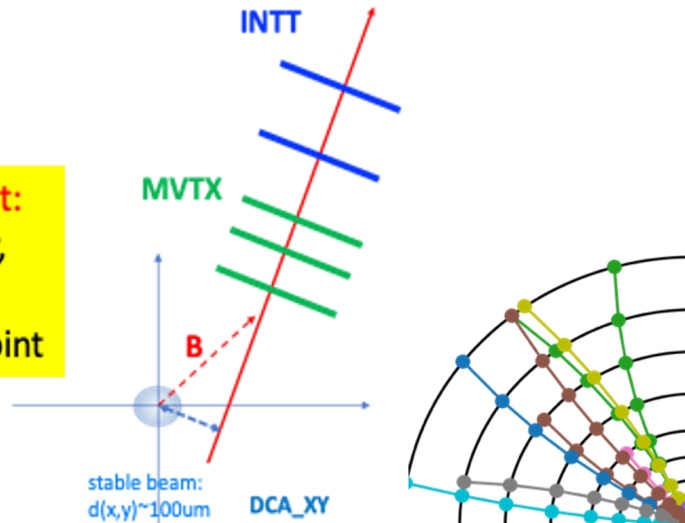
Technical Approaches and Highlights - II

Dantong

- **Objective 2 – Design advanced deep neural networks commensurate with sPHENIX/EIC streaming data requirements.** We aim to design deep neural networks with the following goals: (a) network size: neuron weights that fit in the FPGA block RAMs (BRAMs) of the FELIX cards in sPHENIX/EIC experiments, (b) handling the extremely low signal-to-noise ratio of hit images due to the sparse read-out of the high-resolution MVTX and INTT detectors, (c) performance improvements: 10% improvement over state-of-the-art triggering algorithms, and (d) minimal performance gap between simulated data and real experiment readouts, and outstanding generalization capability. (**Task 2**)

Identify B-hadron event:

- Topology of B decay, with large DCA
- Monitor collision point



Trigger AI Algorithm R&D

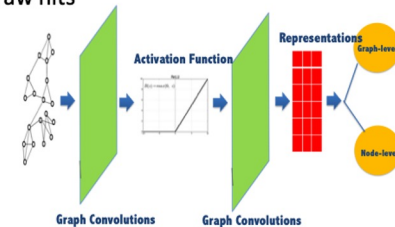
Implemented several models to solve the trigger detection problem:

- Directly applied GNN model to trigger detection problem (GNN)
- Added a global vector to the GNN model to represent some global feature (VPGNN)
- DiffPool model (DiffPool)
- VpGNN + DiffPool (GNNDiffPool)
- ParticleNet, Georgian

Another model we tried: Set2Graph (Affinity Matrix Prediction)

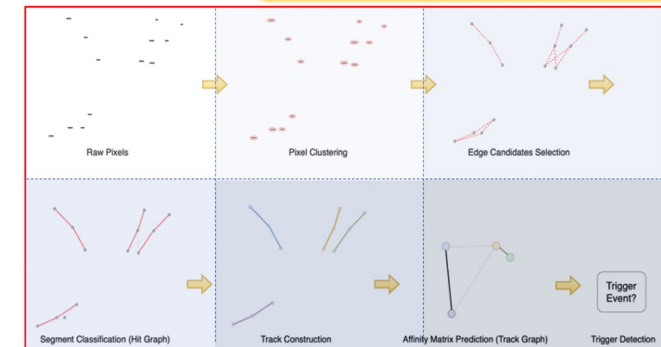
Inputs:

-raw hits



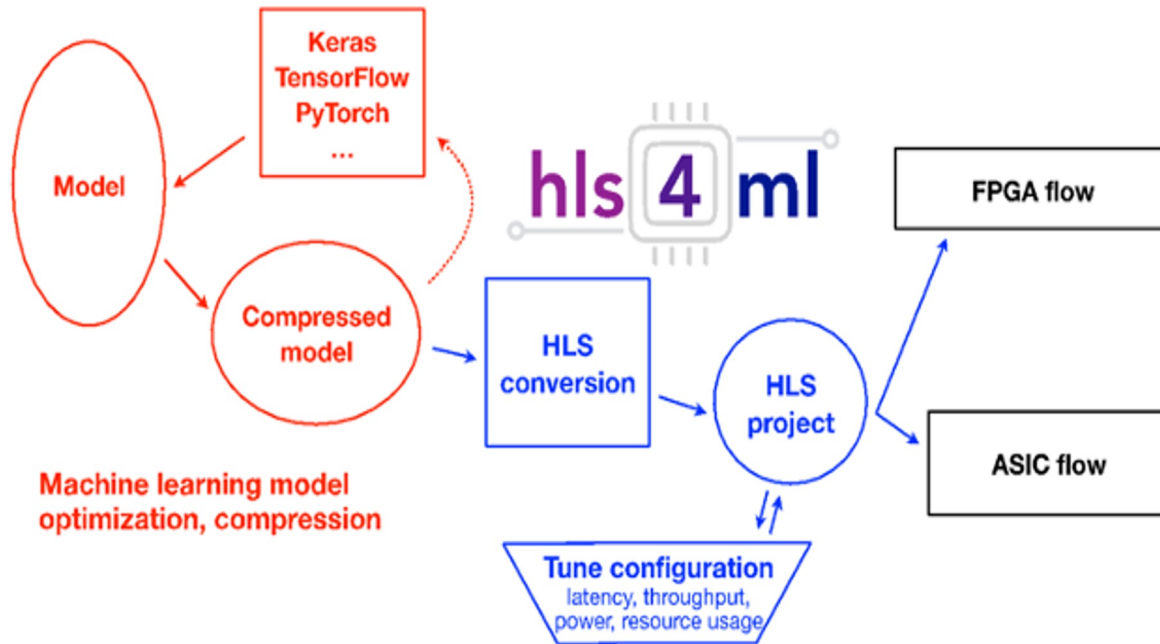
True tracklets:

- 1) 90% BG rejection, Sig_eff ~ 90%
- 2) 99% BG rejection, Sig_eff ~ 40%



Technical Approaches and Highlights - III

- Objective 3 – Deploy advanced deep neural networks within the FELIX system that are capable of real-time reconstruction of heavy flavor events at high throughput. With the development of advanced deep neural networks, a parallel strategy is needed to ensure that these networks can be designed to operate at low latency and high throughput on the FELIX FPGA cards. This challenge involves detailed AI/hardware co-design to ensure that the desired algorithms can be fit within existing resources, and can achieve full throughput. (Task 3)



*Primary focus:
achieving low latency, real-time
processing of data, and
deployment of algorithms with
high efficiency*

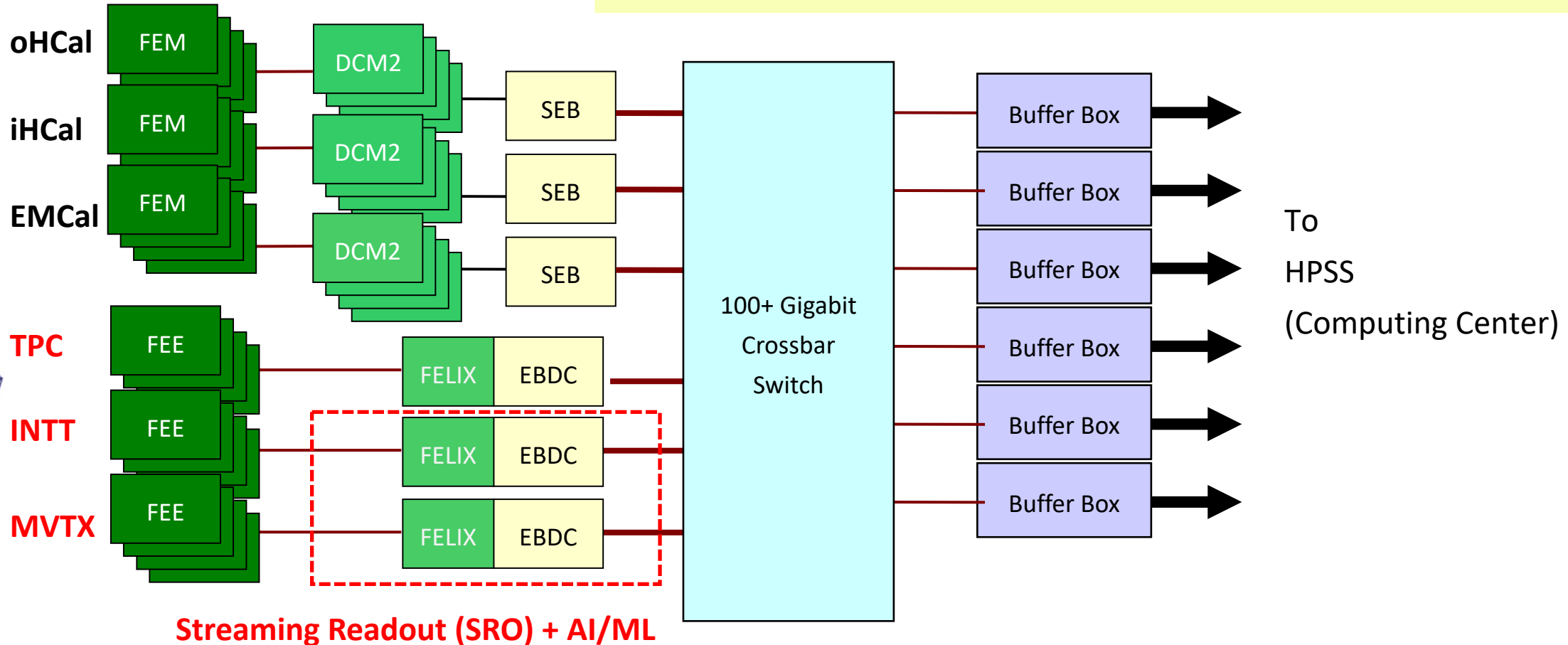
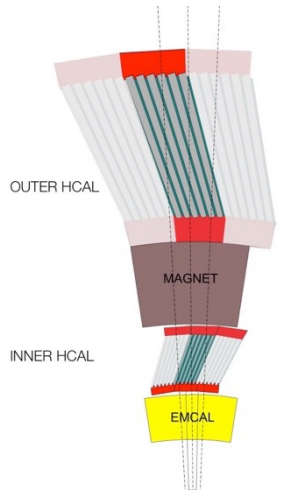


sPHENIX Readout and AI-ML HF Trigger Integration

On Detector

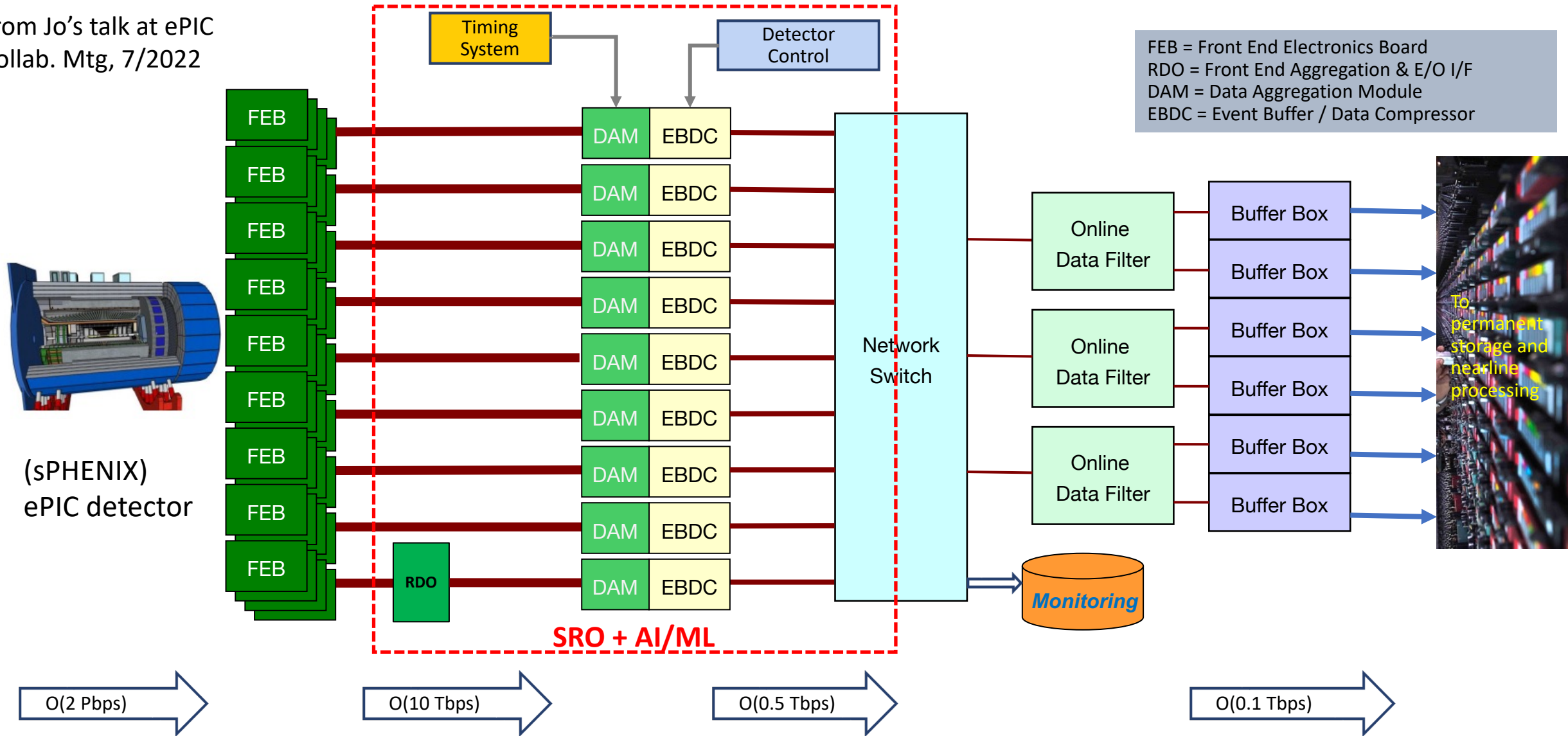
Rack Room

- DCM-2 receives data from digitizer, zero-suppresses and packages
- SEB collects data from a DCM group (~20)
- EBDC Event Buffer and Data Compressor (~40)
- Buffer Box data interim storage before sending to the computing center (6)



From sPHENIX to EIC/ePIC: Streaming + AI/ML DAQ

from Jo's talk at ePIC
collab. Mtg, 7/2022



Summary and Outlook

- Produced full sPHENIX physics and detector simulations of heavy quark and QCD backgrounds
- Successfully developed preliminary AI-algorithms for sPHENIX HF triggers
- Completed MVTX SRO
- Demonstrated INTT SRO
- Implemented a toy AI-algorithm in HLS4ML in FELIX
- Work in progress to implement full sPHENIX trigger in a simplified hardware

On track to accomplish all major goals in 2023

Future plan/proposal:

- Extend project by 2 years, 2024 ~ 2025
 - Implement the demonstrator for sPHENIX p+p run in 2024
 - Develop EIC/ePIC TDR of SRO with AI/ML for EIC CD2(2024) and CD3(2025) based on our work

Table 2: Tasks and Milestones

From DOE proposal

Note: completed; in progress

TASK I	TASK II	TASK III	TASK IV
Year 1			
1. Software development and system design. We will first perform detailed sPHENIX physics and detector simulations to design a real-time fast data processing and autonomous detector control and calibration system. In the meantime, we will survey currently available AI models and design a system for offline training and domain adaptation for data and MC. The physics and detector simulation results and the performance of hardware are used to tune the AI algorithms. 2. Hardware development and system integration. We will take advantage of the streaming readout capability of the sPHENIX tracking system to implement continuous readout of two fast silicon tracking subsystems, MVTX and INTT. A FPGA based fast tier-1 AI system will be developed to identify heavy flavor (HF) events in $p+p$ collisions, and generate a fast trigger to initiate the readout of TPC.			
By Q2 ► Generate open heavy flavor and QCD background events for simulations (LANL, MIT)	► HF trigger algorithm development for FPGA (FNAL, LANL, NJIT)	► MVTX streaming read-out (LANL) ► INTT streaming read-out (MIT)	► Beamspot interaction and readout simulation (FNAL) ► Displaced tracks and anomaly simulation (FNAL, MIT)
By Q3 ► Develop fast tracking algorithms using MVTX and INTT hit information (LANL, MIT, NJIT)	► Design real-time GPU training machine (MIT, NJIT)	► hls4ml implementation and algorithm development (FNAL, MIT)	► Preliminary design of streaming and automated controls of online GPU-based training system (MIT, NJIT)
By Q4 ► Complete a preliminary design of HF trigger AI offline (FNAL, LANL, NJIT)	► ML, Graph NN training, by NJIT and MIT	► FPGA implementation of HF trigger with MVTX and INTT (All)	► Simulation and training (MIT, NJIT)
Year 2			
We will focus on the system integration and continue to improve and benchmark the performance of software, firmware, and hardware system.			
By Q5 & Q6 ► Interface between AI system and MVTX detector Data Input by (FNAL, LANL) ► Interface between AI system and TPC Readout Control (FNAL, LANL)	► Design new GNNs (Encoder, Attention, Particle-Net) algorithms with hls4ml (MIT, NJIT)	► hls4ml customization for FELIX board (FNAL, LANL) ► FPGA, GPU system integration and evaluation (MIT, NJIT)	► GPU deployment for autoencoder and training (FNAL, MIT)
By Q7 ► Continue to improve algorithms for HF tagging (LANL)	► Improve algorithm with hls4ml on FPGA (FNAL, NJIT)	► Multi-FELIX Board Integration (LANL) ► Validation and test with FELIX boards (All)	► ML model and domain adaption update (MIT, NJIT)
By Q8	11/11/22 ► Benchmark system performance with sPHENIX or test beam data (All)		

Physics simulation and AI-ML algorithms

- Dantong Yu



AI Trigger Pipeline

1. Fetch events from event buffer to Processing



2. Data Pre-processing Clustering



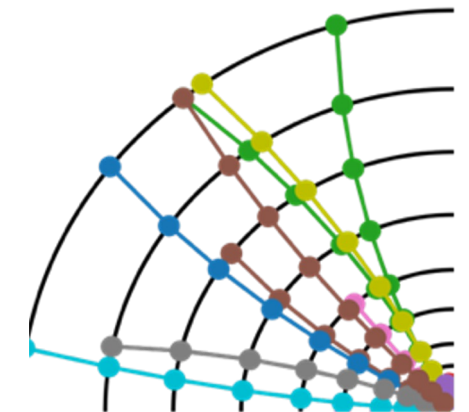
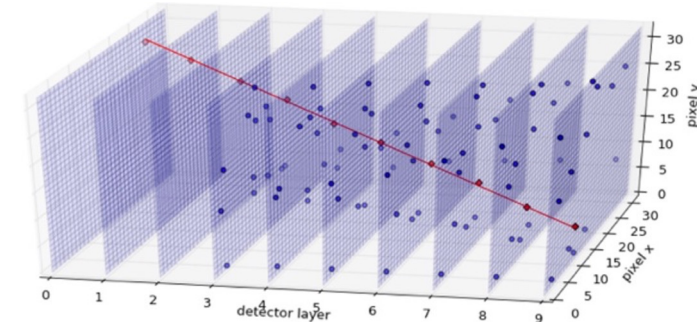
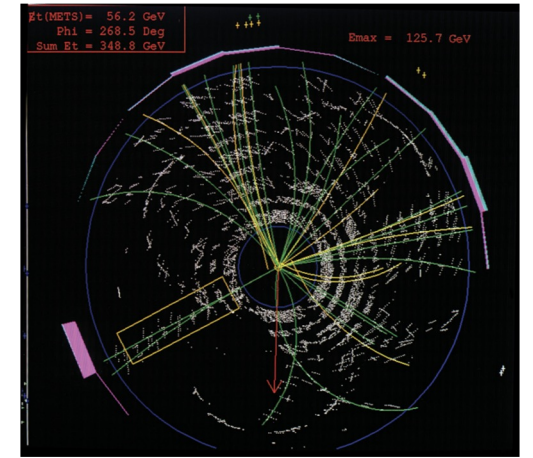
3. Tracking + Outlier hits Removal



4. Triggering

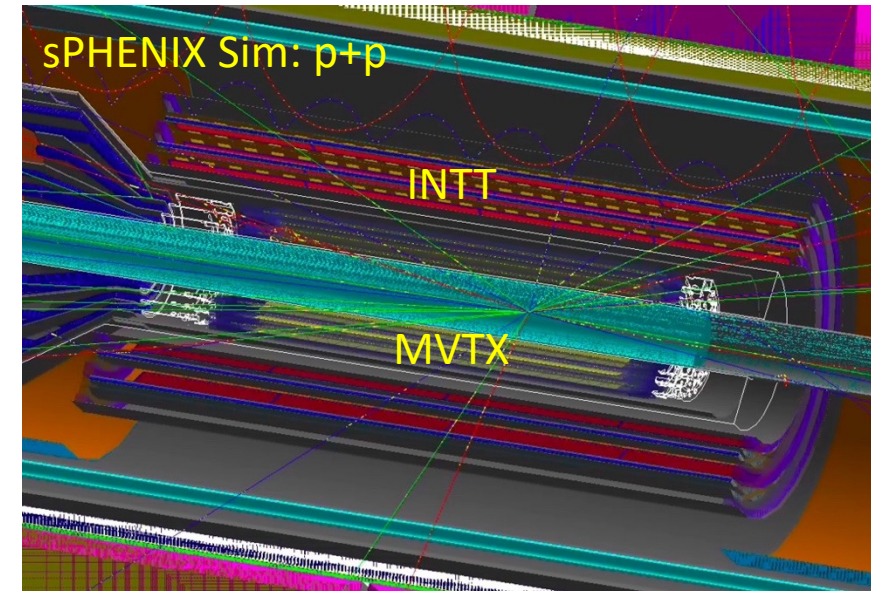


5. Triggers on TPC



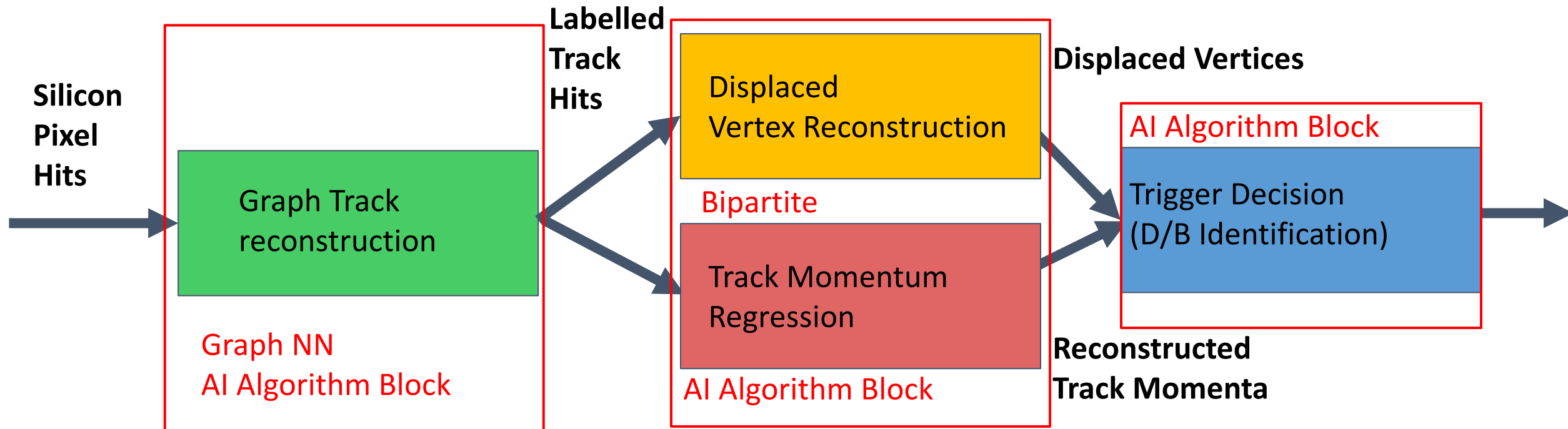
Simulations

- **Most simulations have focused on sPHENIX**
 - ePIC detector simulations for EIC are under rapid changes
- **We simulated 2 different physics events (all 200 GeV pp collisions)**
 - Minimum bias with all heavy flavor decays rejected for background
 - $D^0 \rightarrow K^- \pi^+$ in minimum bias for signal
- **Simulations have been improved throughout the year**
 - Added full service material for MVTX
 - Added realistic hit duplication in MVTX from pixel pulse length
- **Dataset**
 - 8 million events - 50% signal/noise.
 - 8 million events - 0.1% signal/noise.

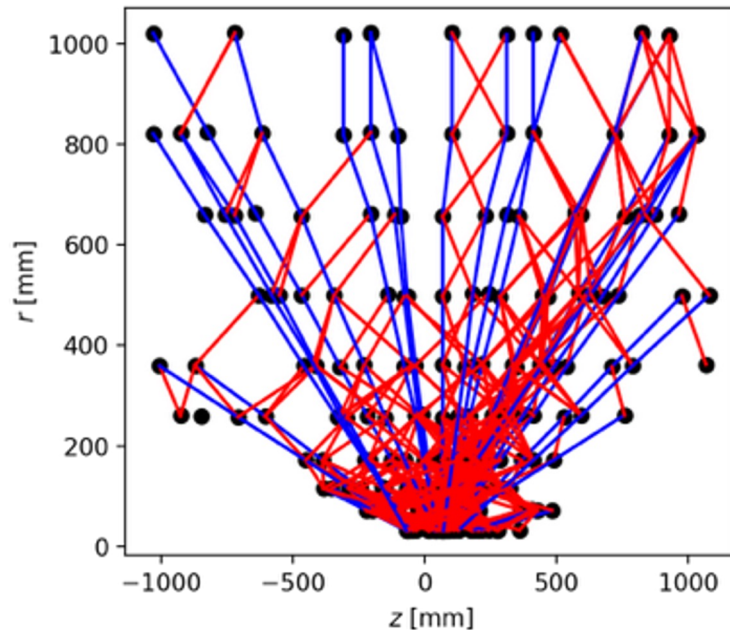


Algorithm Flow

- **Strategy is to construct a large scale AI algorithm with all elements of trigger**
 - Algorithm is factorizable into core physics components
 - Emulates the normal reconstruction workflow
 - Use of core components allows for intermediate physics validation



Tracking as Graph Neural Networks



- **Graph neural networks are a popular way of posing tracking problems in physics**
 - Sparse hits in space do not fit traditional ML architectures
 - Growing sub-field of geometric deep learning
- **Data is structured as a graph of connected hits**
 - Connect plausibly related hits using geometric constraints and learn best associations and parameters of connected hits (tracks)

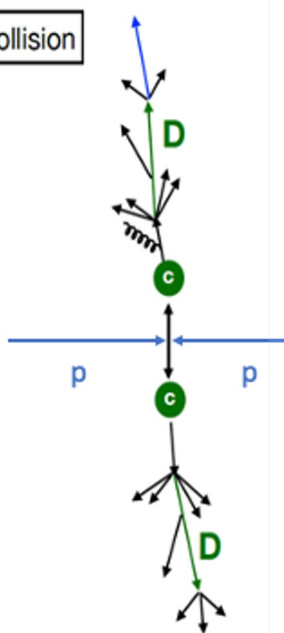
Three Types of Potential Interactions in Model Event Decay

Local track-to-track interactions, such as determining whether two tracks share the same origin vertex.

Track-to-global interactions, such as determining the collision vertex of the event and potential secondary vertices of decaying.

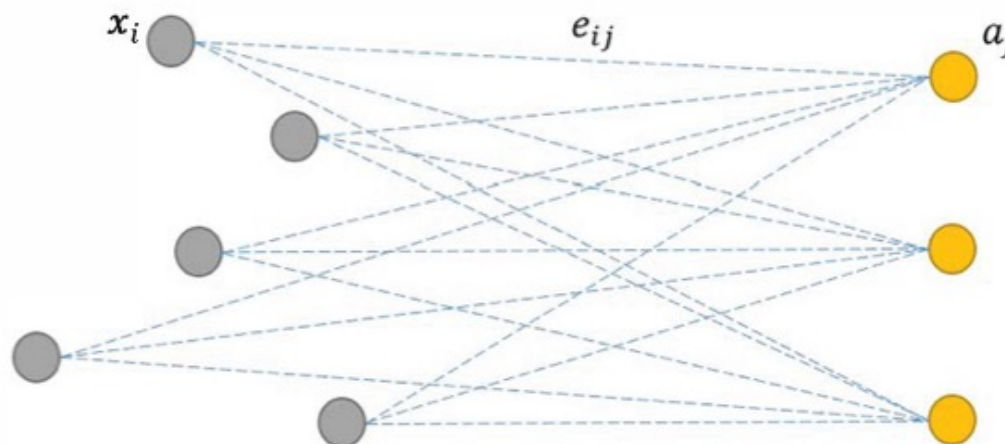
Global-to-track interactions, such as comparing each track's origin with the collision vertex of the event.

pp collision



Track Nodes

Vertices



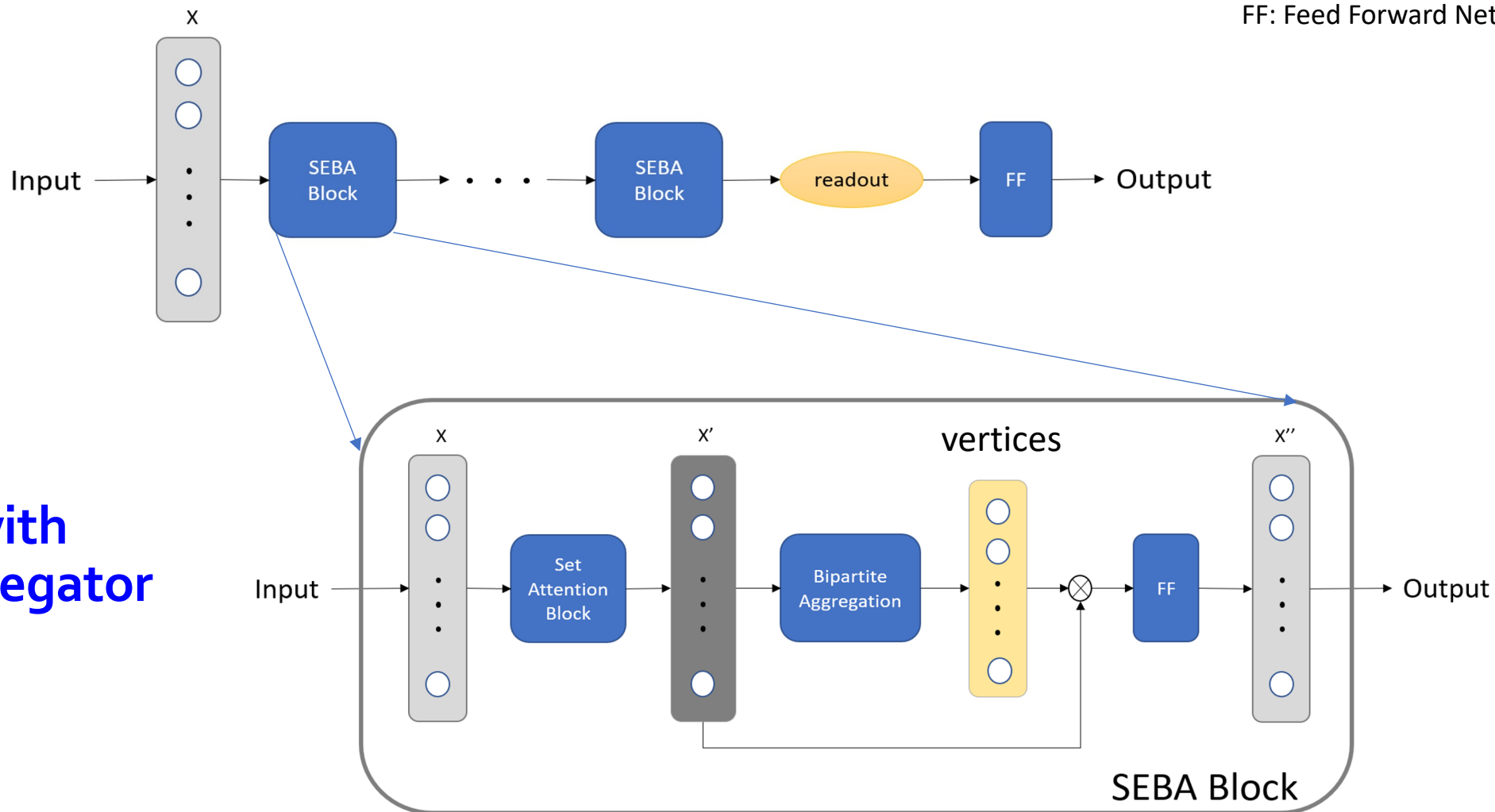
**Pair-wise Relations
(Set/Graph/Attention)**

**Graph Aggregation
(Bipartite Graph aggregation)**
- vertexing

**Graph Distribution and
Bypassing (Bipartite Graph)**
- parent/daughter

Bipartite Graph Networks with Set Transformer (BGN-ST) Model Architectures

FF: Feed Forward Network



**Set Encoder with
Bipartite Aggregator
(SEBA) Blocks**

Experiments: Physics Driven BGN-ST

Model	#Parameters	Accuracy	AUC
Set Transformer	80,002	86.17%	91.75%
GarNet	284,210	86.22%	91.81%
PN+SAGPool	780,934	86.25%	92.91%
BGN-ST	363,426	87.56%	93.22%

- Excellent physics trigger performance for charm events
- Expect better results for beauty, work in progress

1% signal/background ratio		
Background Rejection	Efficiency	Purity
90%	72.5%	7.25%
95%	48.9%	9.78%
99%	15.0%	15.0%
99.33%	10.5%	15.7%

Status and Outlook

- **Demonstrated BGN-ST outperforms selected state-of-the-art methods** ^[1,2]
 - Adopts the physics-aware concept and introduces explicit physics properties such as transverse momentum
 - Improves the task accuracy and AUC score by about 15%
 - Architecture benefits from pairwise interactions between tracks and allows a two-way scattering and gathering for effective information exchange and adaptive graph pooling
- **Next step: from optimal performance of state-of-the-art methods, develop firmware implementations**
 - Understand what is feasible and condense the model into latency and resource constraints

[1] Xuan, Y. Zhu, G. Borca-Tasciuc, M. X. Liu, Y. Sun, C. Dean, Y. C. Morales, Z. Shi, D. Yu, End-To-End Pipeline for Trigger Detection on Hit and Track Graphs in, Accepted by Thirty-Fifth Conference on Innovative Applications of Artificial Intelligence, IAAI.

[2] Xuan, G. Borca-Tasciuc, Y. Zhu, Y. Sun, C. Dean, Z. Shi, D. Yu, Trigger Detection for the sPHENIX Experiment via Bipartite Graph Networks with Set Transformer in European Conference on Machine Learning and Principles and Practice of Knowledge Discovery in Databases (ECML-PKDD 22).

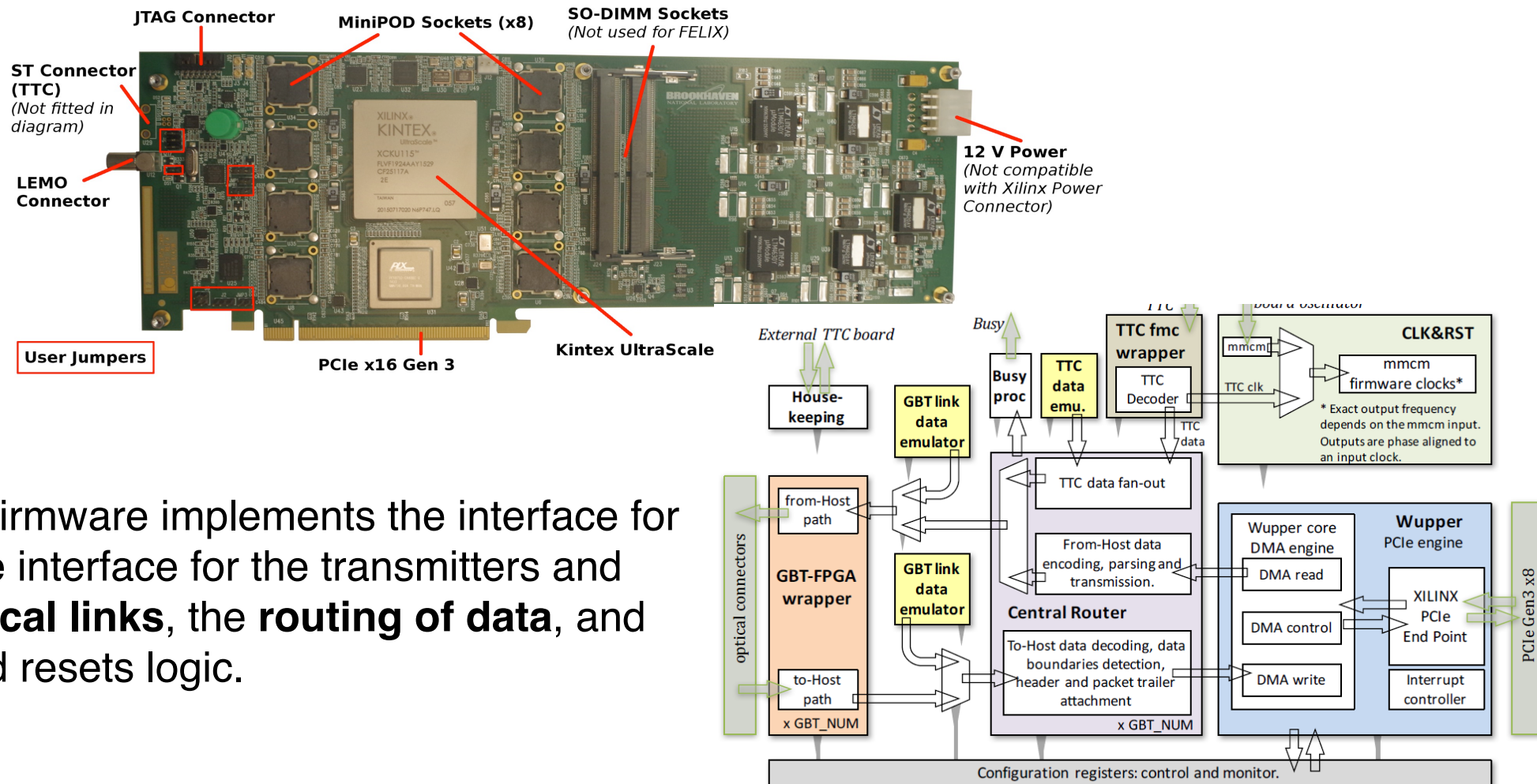


hls4ml translation and firmware implementation

- Micol Rigatti and Phil Harris

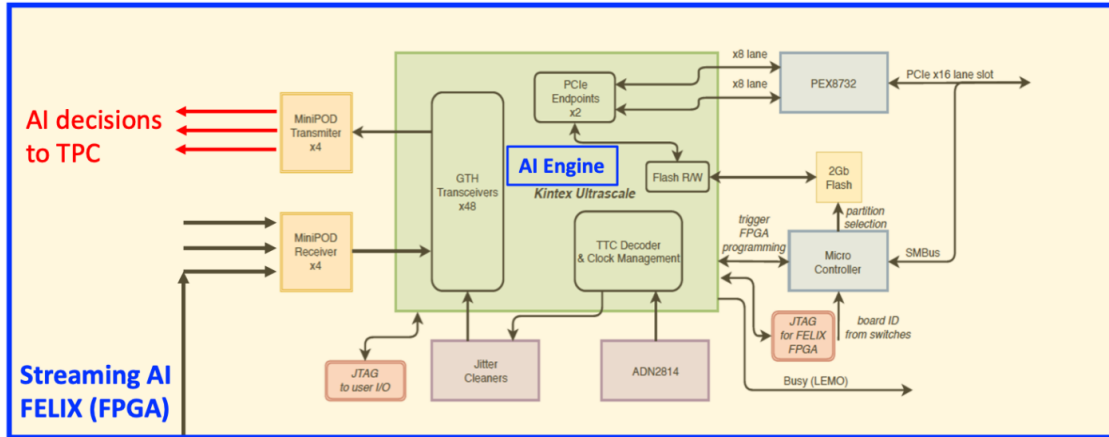
Hardware and Firmware Implementations

Selective data streaming in sPHENIX will use the FELIX board. FELIX is a 16-lane Gen-3 PCIe card with 48 transmitters and receiver optical links. The on-board FPGA is a Kintex Ultrascale XCKU115FLVF1924-2E

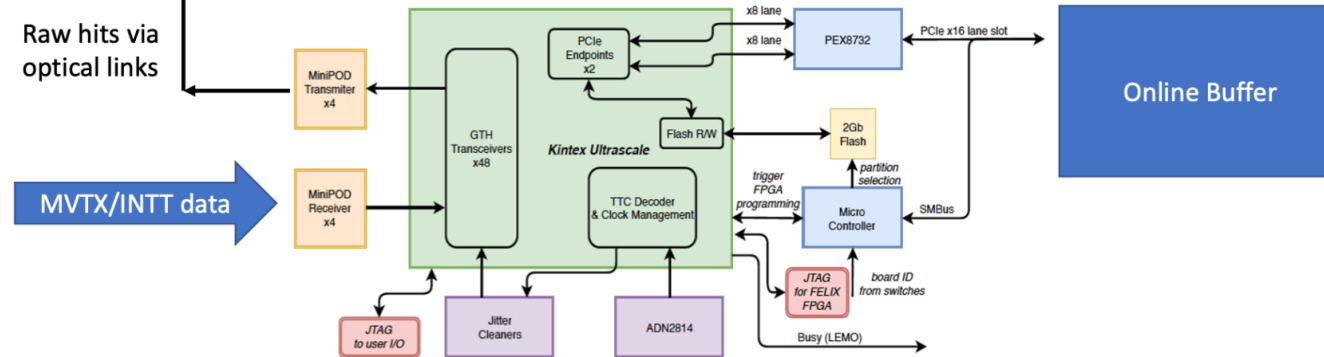


The FELIX Firmware implements the interface for the **PCIe**, the interface for the transmitters and receiver **optical links**, the **routing of data**, and the clock and resets logic.

Hardware Configuration



A second FELIX is connected to the first FELIX through the optical transceivers, as a **dedicated FPGA hardware for smart control and real-time decision-making** for TPC readout in the selective data streaming architecture.



The **AI Engine** will search for displaced tracks to identify tracks from **heavy quark decays** that are pointing away from the nominal beam center. Such a signal will initiate the readout of the TPC detector.

The firmware of the AI Engine will implement the **Neural Network** deployed using **hls4ml**, the PCIe interface, and the optical link interface.

Schedule – completed

1

The Fermilab **test stand** has been set up: the FELIX board **BNL711** is set up in a Host Computer.



2

On the Host Computer, the **FELIX driver** and the **software application** is been installed.

3

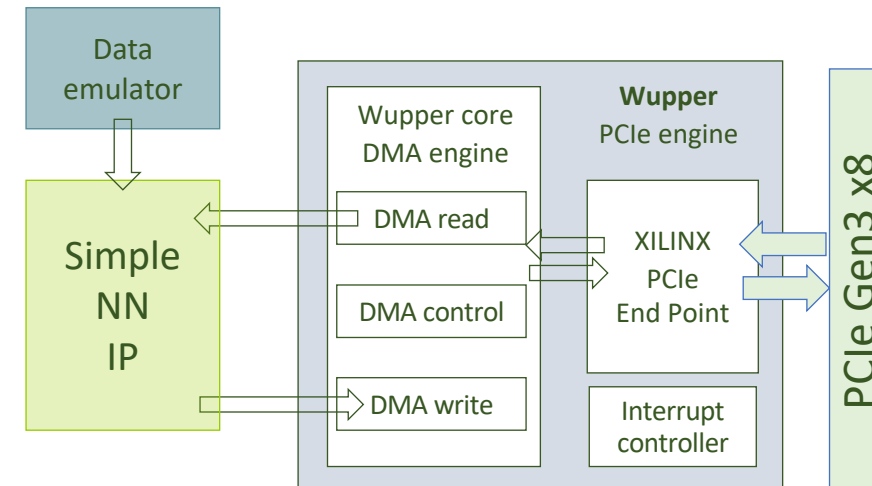
The **data movement** between the BNL711 and the Host Computer is been tested. The data are emulated inside the BNL711.

4

The hsl4ml workflow is been tested with the Kintex Ultrascale XCKU115 as a target generating the IP of a **simple NN**

5

The IP of the simple NN is in the integration phase with the **Wupper** module (**PCIe engine**)



Schedule – next steps and challenges

1

Test the functionality of the NN and the data routing mechanism



The FELIX card is a **data router**. It needs an application to be instructed on the data movement.

The **application** relies on an **OS** and a **driver** interfacing the FELIX card.

- Inserting the NN in the FELIX Firmware will change how the driver interfaces with the card
- The testing part right now relies on the interplay of Hardware, Firmware, and Software



Interfacing the board to test the NN is the other real challenge

2

Fit the NN in the FELIX Firmware complying with timing constraints.

The workflow of hls4ml generates an IP considering available the resources of the target FPGA. We need to route the NN considering the already routed FELIX Firmware.



Meeting the timing in this condition is the real challenge

3

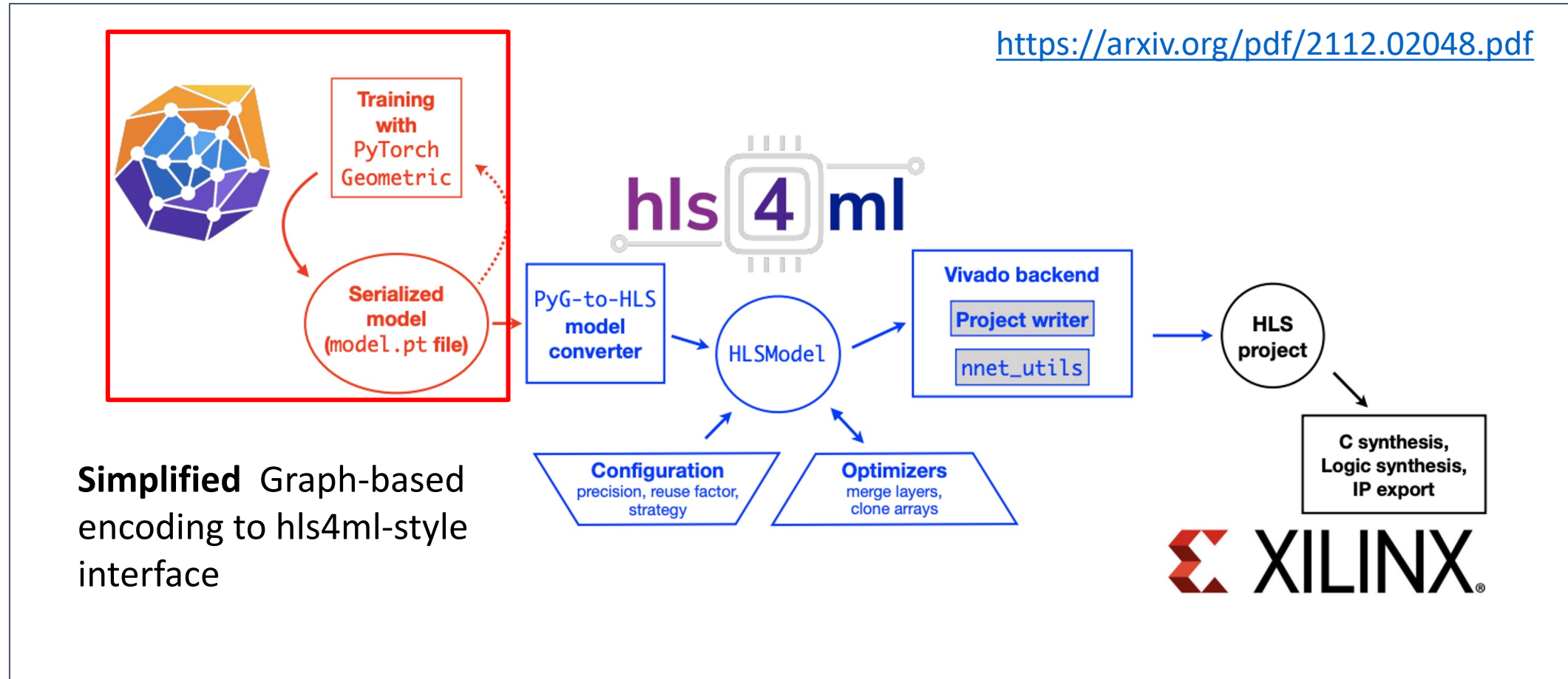
NEXT STEP: the analysis of the FELIX Firmware in terms of resources to drive the hls4ml implementation of the NN

4

NEXT STEP: the integration of the IP of the simple NN is in with the Wupper module (PCIe engine)

GNN Implementation in hls4ml

First version of hls4ml-based Graph encoder available



Critical optimization needed in construction of Graph-based mapping

GNN Planned Upgrade in hls4ml

Design	$(n_{\text{nodes}}, n_{\text{edges}})$	RF	Precision	Latency [cycles]	II [cycles]	DSP [%]	LUT [%]	FF [%]	BRAM [%]
Throughput-opt.	(28, 56)	1	ap_fixed<14, 7>	59	1	99.9	66.0	11.7	0.7
Throughput-opt.	(28, 56)	8	ap_fixed<14, 7>	75	8	21.9	23.8	4.7	0.7
Resource-opt.	(28, 56)	1	ap_fixed<14, 7>	79	28	56.6	17.6	3.9	13.1
Resource-opt.	(448, 896)	1	ap_fixed<14, 7>	470	174	56.6	25.0	7.4	16.5
Resource-opt.	(448, 896)	8	ap_fixed<14, 7>	1590	520	5.6	25.0	7.4	16.3

@200 MHz, 1590 Cycles $\rightarrow$ 7.5 μ s

Active work is underway to improve the GNN implementation

- Base implementation has been updated twice in 2022
 - A third update will start in mid December, focusing on 3 examples
 - Example 1: Tri-muon reconstruction with the LHC (muon endcaps)
 - Example 2: Heavy flavor tracking at sPHENIX
 - Example 3: Silicon strip tracking at LHC
 - Current Graph encoder to be optimized (adjacency matrix computation)
 - Aiming to centrally rework this with core hls4ml developers
 - Will be a central project across several domains



Demonstrator Implementation

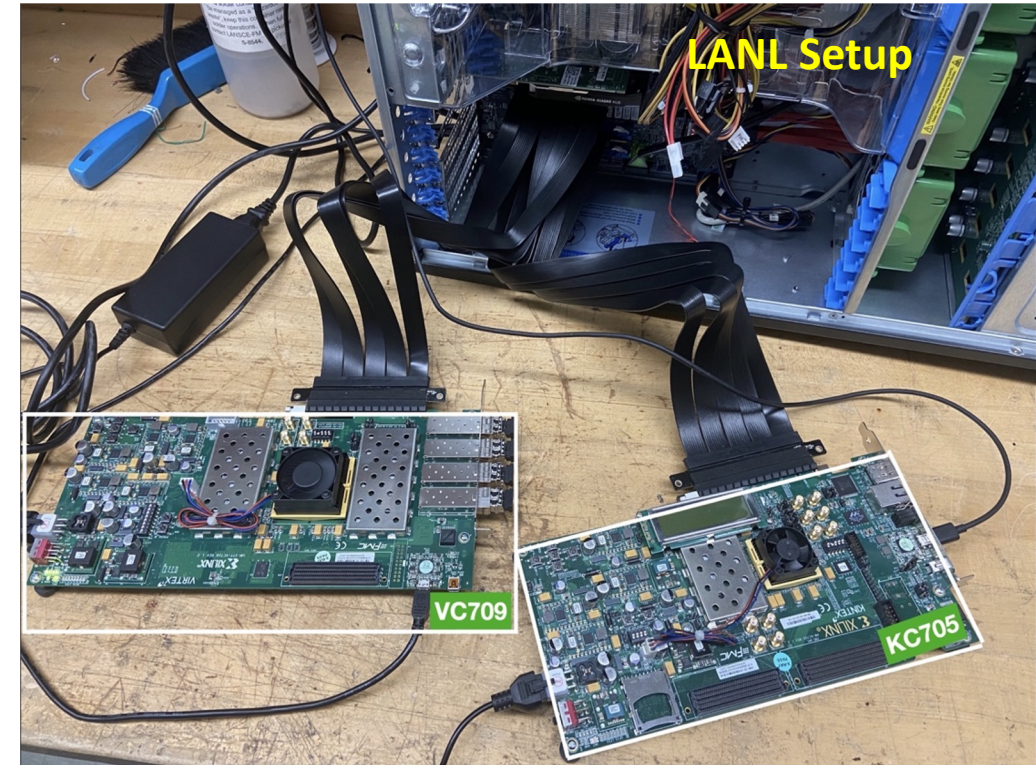
- Jakub Kvapil

Demonstrator Development @ LANL, ORNL and CCNU

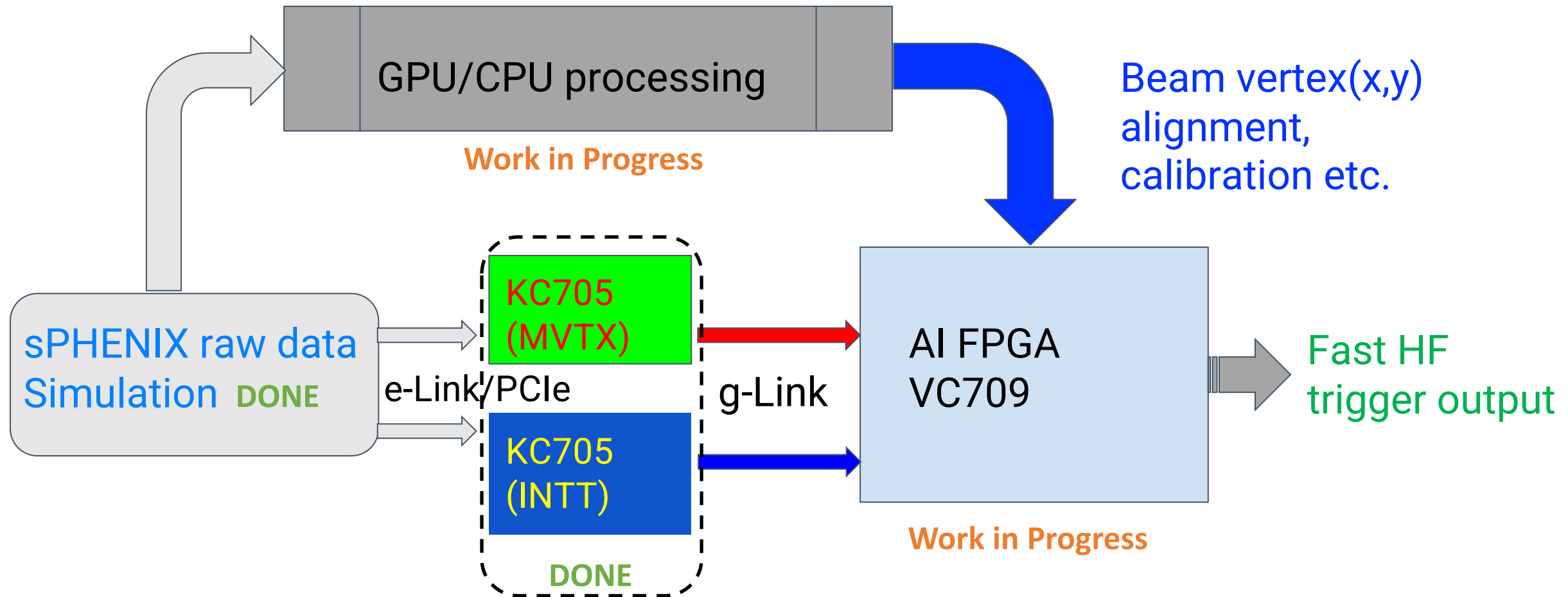
Why a demonstrator?

- sPHENIX technologies rapidly evolving at the time of proposal

1. sPHENIX DAQ will be ready in early 2023
 - Parallelize development in early stage to be ready for final deployment
2. Fabricate more FELIX boards later
 - The AI core logic will be implemented on VC709 which is the FELIX protoboard (share similar resources and is supported)



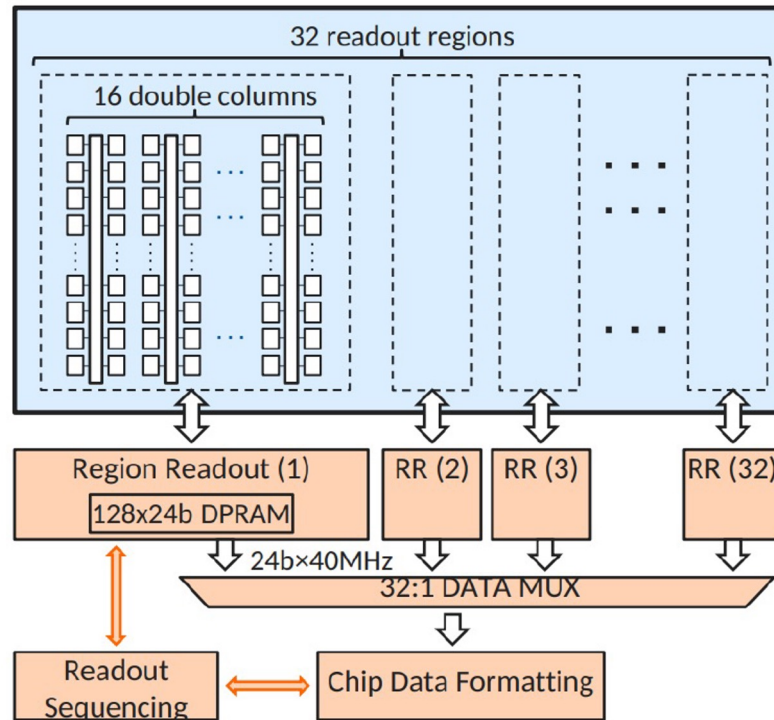
Demonstrator Implementation



Both KC705 and VC709 will be replaced by FLX712 at deployment stage

sPHENIX RAW Hit Data Simulation

- In order to test the full loop feedback AI system detector hits must be known
- Code developed to transform MC JSON hit patterns into MVTX and INTT raw data streams

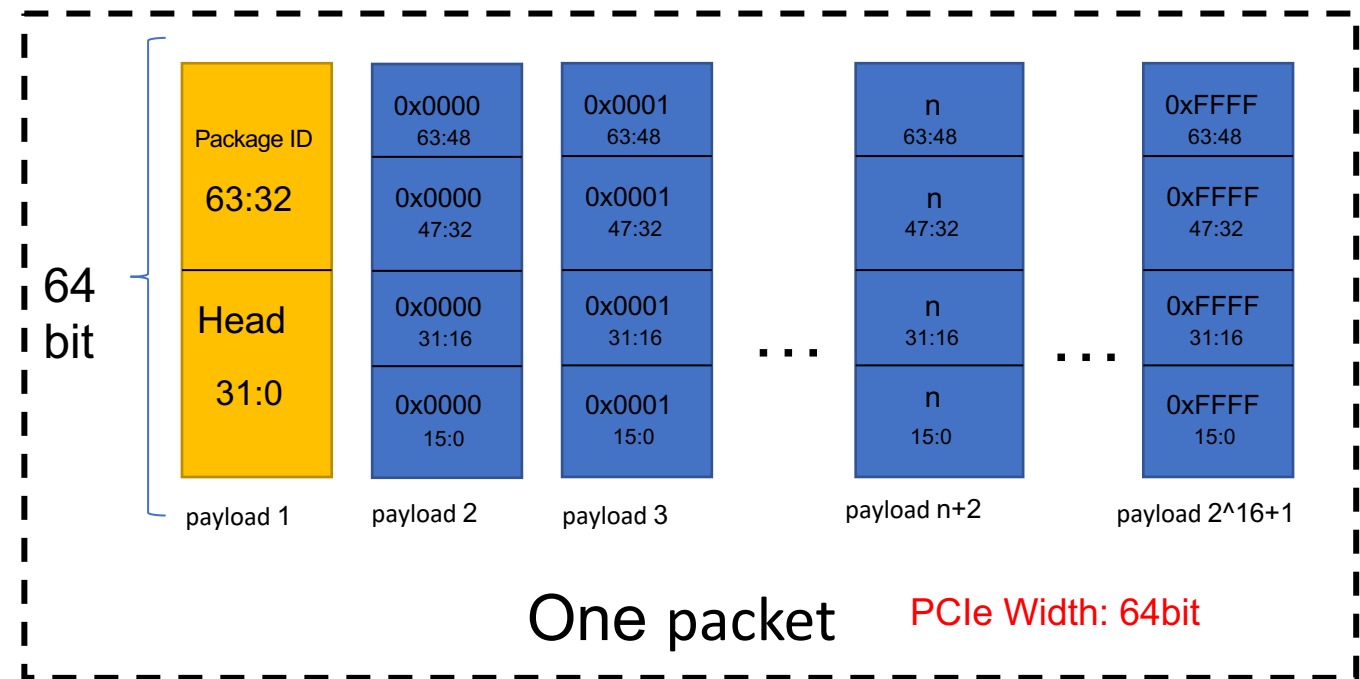
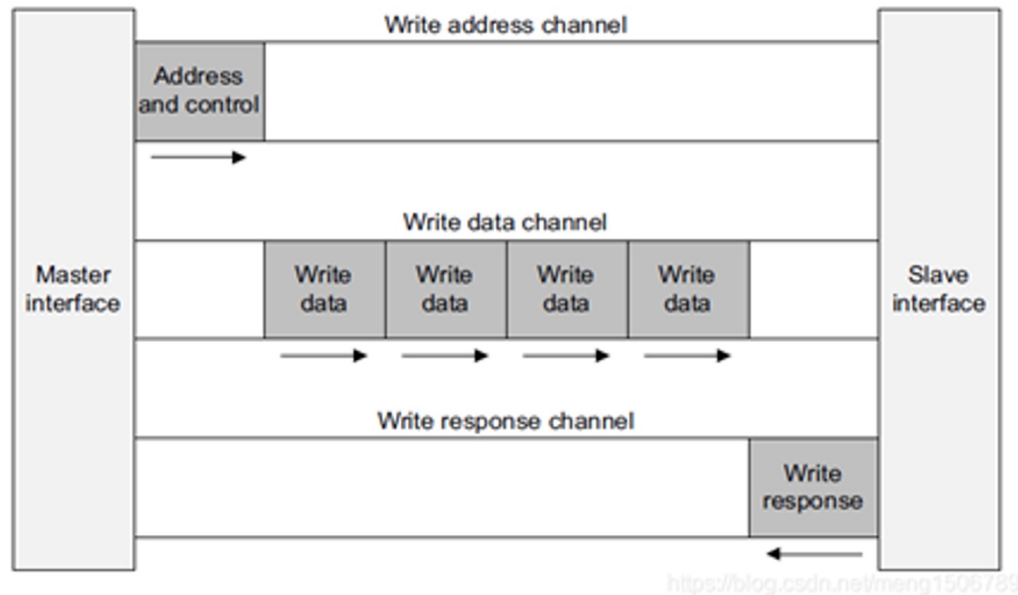


MVTX sensor readout

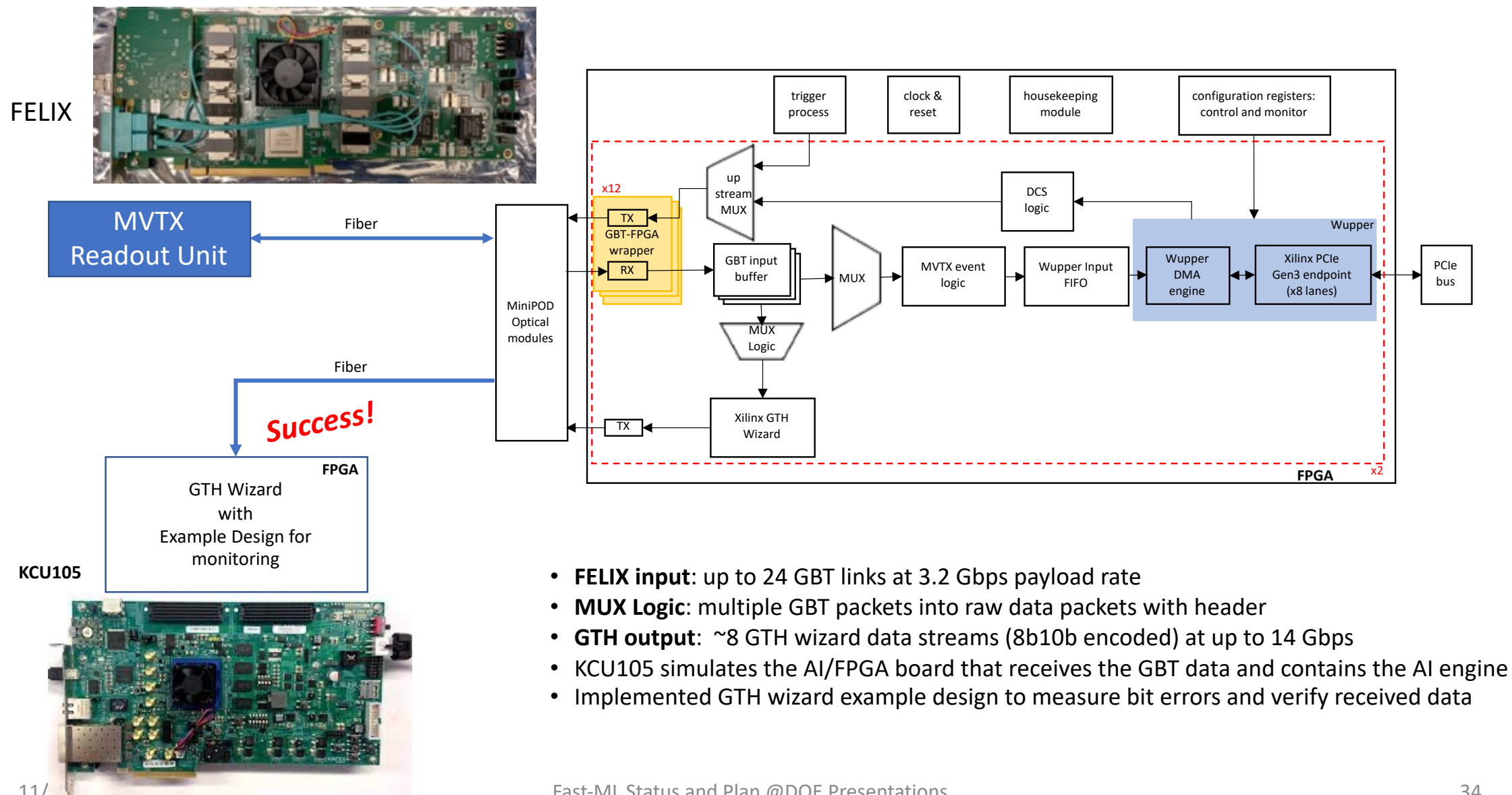
```
"MVTXHits": [
  {
    "ID": {
      "HitSequenceInEvent": 0,
      "G4HitAssoc": [
        20
      ],
      "MVTXTrkID": 31,
      "Layer": 0,
      "Stave": 1,
      "Chip": 6,
      "Pixel_x": 260,
      "Pixel_z": 72
    },
    "Coordinate": [
      1.7037016547341806,
      1.8503514306175895,
      4.744902043342591
    ]
  },
  ]
```

KC705 MC Raw Data Transmission Demonstrated

- Simulated data must be streamed to the AI/FPGA, VC709/FLX712
- KC705 is used to read the MC data via PCI XDMA and transmit them using g-Links
- One link for MVTX and one for INTT



FLX712 Raw MVTX Data Transmission to AI-FPGA



Timeline and Outlook

2021

2022

2023

2024

2025

2030+



• Project funded by DOE FOA - Dec. 2021 (Lab) - Jan. 2022 (Univ.)	• MVTX & INTT SRO • Fast tracking & trigger algorithms in place • Initial FPGA bitstream • GPU feedback machine R&D	• Refine interface between system and detectors • Improve algorithms with latest data stream • Pre-commissioning	• Deploy device at sPHENIX pp/pAu run • EIC preliminary TDR (CD2)	• Final design for EIC TDR (CD3) • Take advantage of new technology if required	• Deploy device at EIC
---	--	--	--	--	--

Great progress, on track to succeed in 2023

Future plan, 2024+

Backup slides

sPHENIX and EIC Schedules

sPHENIX Run Plan: 2023-2025

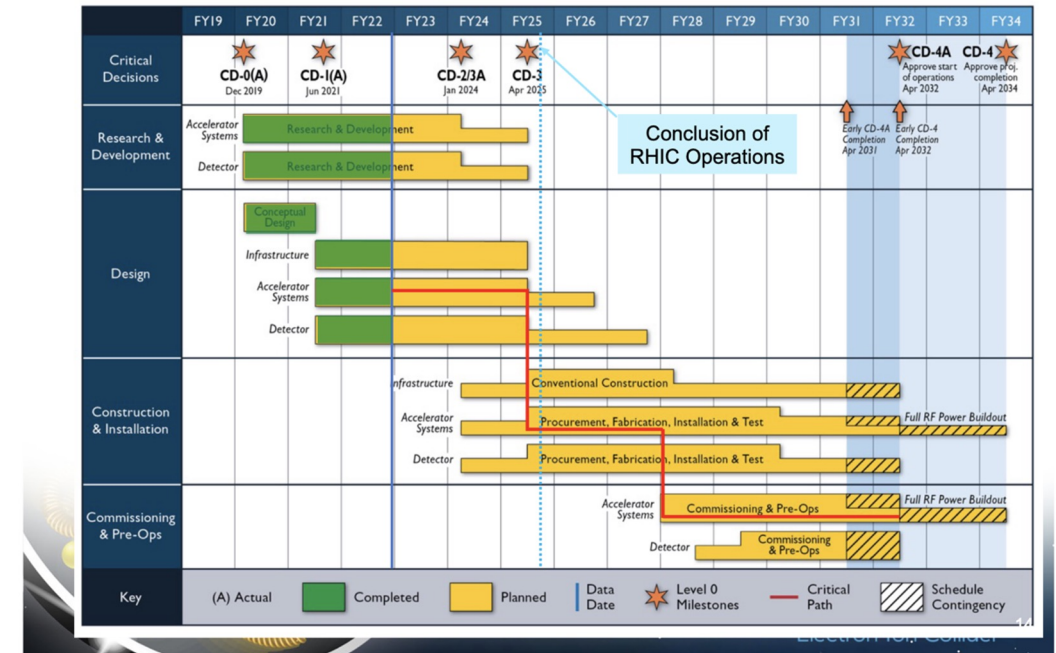
- pp and pAu run in 2024

Year	Species	$\sqrt{s_{NN}}$ [GeV]	Cryo Weeks	Physics Weeks	Rec. Lum. $ z < 10$ cm	Samp. Lum. $ z < 10$ cm
2023	Au+Au	200	24 (28)	9 (13)	3.7 (5.7) nb ⁻¹	4.5 (6.9) nb ⁻¹
2024	$p^\uparrow p^\uparrow$	200	24 (28)	12 (16)	0.3 (0.4) pb ⁻¹ [5 kHz] 4.5 (6.2) pb ⁻¹ [10%-str]	45 (62) pb ⁻¹
2024	$p^\uparrow$ +Au	200	–	5	0.003 pb ⁻¹ [5 kHz] 0.01 pb ⁻¹ [10%-str]	0.11 pb ⁻¹
2025	Au+Au	200	24 (28)	20.5 (24.5)	13 (15) nb ⁻¹	21 (25) nb ⁻¹

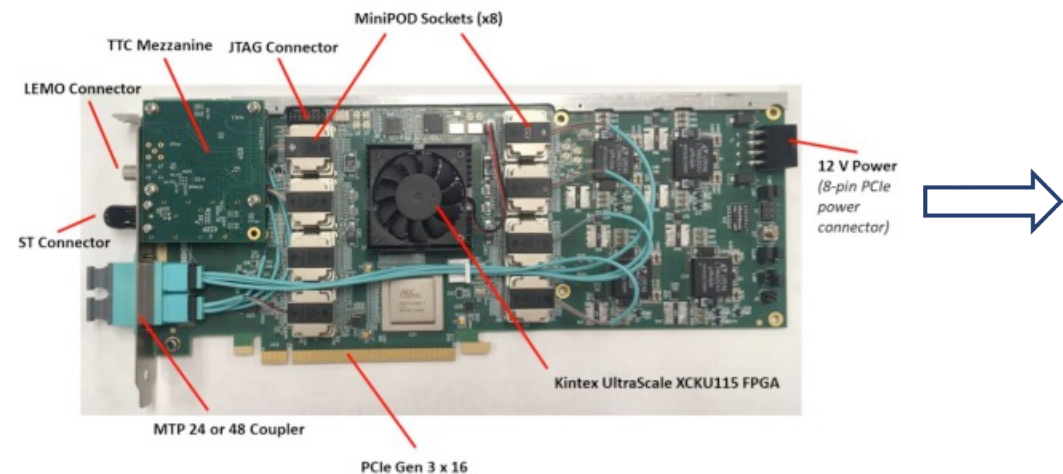
EIC Project Plan (as of 11/05/2022)

- ePIC TDR for CD2/CD3A, 2024
- Final design/construction, 2025

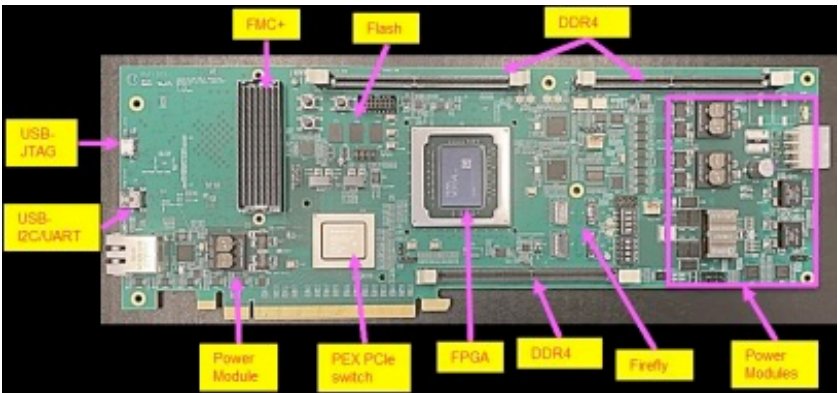
EIC Reference Schedule - V3



EIC: ePIC DAM Candidates: ATLAS “FELIX”



Current ATLAS Phase 1/ sPHENIX FELIX BNL-712v2



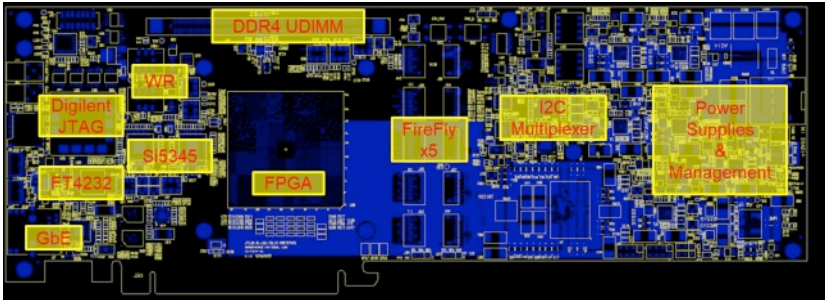
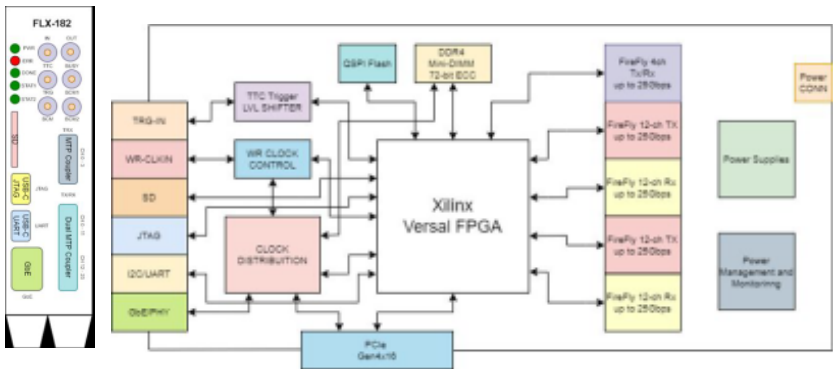
FELIX FLX-181 Prototype (BNL)

(Hao Xu, BNL, DAQ WG meeting 7/2021)



Assembled 24ch FLX-181 with 25 Gbps FireFly FMC+

Current Design of FELIX FLX-182



FELIX Status

FLX-182 Status

- Design passed FELIX review, will be sent out for fabrication in this week
- First assembled board is expected to be delivered in early September 2022
- 7 boards will be produced if there's no big design issues, by December 2022
- Small production for more boards is possible once FPGA is available

Plan for 48-ch FELIX

- FPGA: Versal Premium, e.g. VP1552
- Transceivers: Up to 100+ GTYP/GTM
- PCIe Gen 5 up to 16 lanes
- If FPGA is available as planned, design will start in Q1 of 2023, first board is expected to be available in Q3 2023.

Architecture and Interfaces

- PCIe Gen 4 x 16 lanes
- Transceiver
 - Transceiver Type: Samtec FireFly transceiver
 - Transceiver Speed: up to 10 Gb/s ("CERN-B") or 25 Gb/s
- Number of Optical Connectors per Card
 - At least 24 bi-directional connections to front-end electronics
 - A separate bi-directional connection to the TTC/BUSY system
- Configuration
 - Boot from JTAG/QSPI/SD card
 - Remote FPGA configuration from Multiple Flash Partitions
- DDR4/Flash Memory/SD card
- I2C
- External Electrical Interface
- Voltage Protection
- Temperature Protection