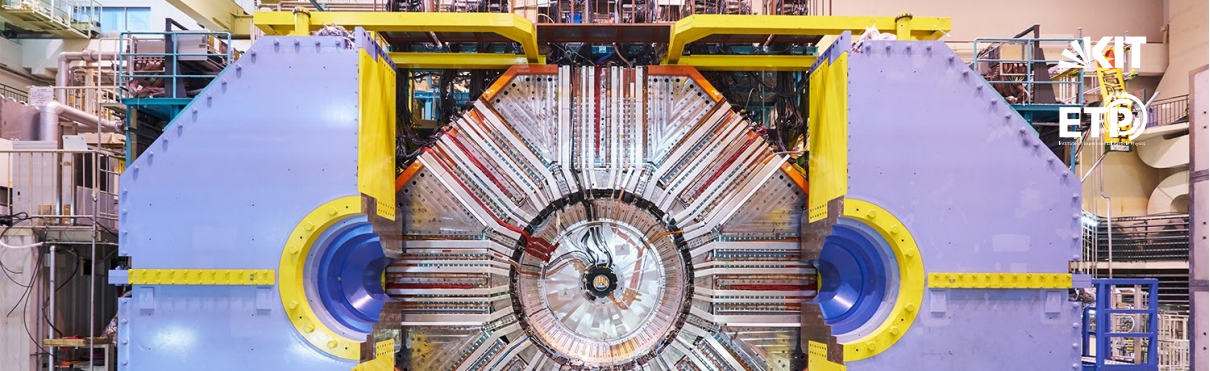




ECLTRG Logic Upgrade

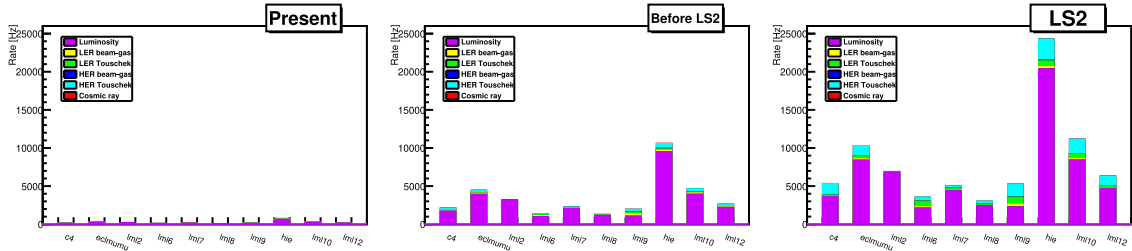
TRG-DAQ Workshop, 23. October 2025

Isabel Haide* (isabel.haide@kit.edu) with input from (alphabetically): Giacomo De Pietro, Patrick Ecker, Torben Ferber, Thomas Lobmaier, Marc Neu, Fabio Papagno

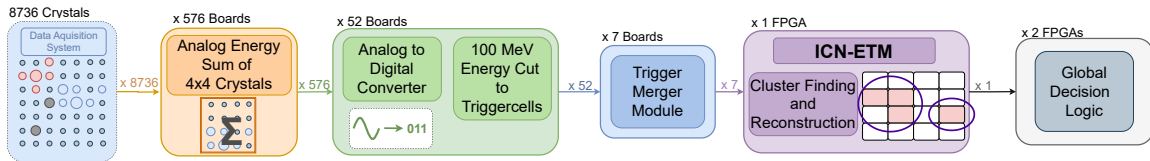
* Institute of Experimental Particle Physics (ETP)

Trigger Rates in Future Conditions

- Trigger rate extrapolations based on dedicated beam time studies ([Belle II Note](#)) show significantly increased rates of ECL trigger bits with higher beam currents and higher luminosities. Luminosity background is the highest contribution for the ECL trigger.
- Loose trigger bits such as `hie` are used by many analyses and will have to be prescaled in future conditions to keep within the DAQ limit.
- Improved readout strategies and new clustering algorithms have to be employed to keep a high trigger efficiency for all analyses.



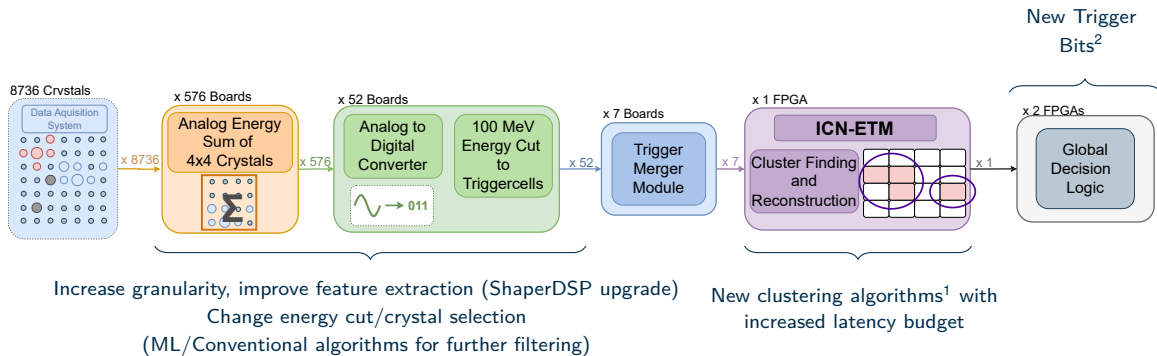
The Current ECL Trigger: ICN-ETM



Current Challenges:

- Coarse granularity of input results in limited position resolution, no resolution for overlapping clusters.
- Coarse granularity of input does not preserve shower shape information.
- ICN-ETM can only return up to 6 clusters per event.
- Loose trigger bits have to be prescaled to reduce overall trigger rate with higher luminosity, leading to a loss in efficiency for analyses. Trigger bits such as c2 (two clusters or more in the inner ECL) are already turned off.

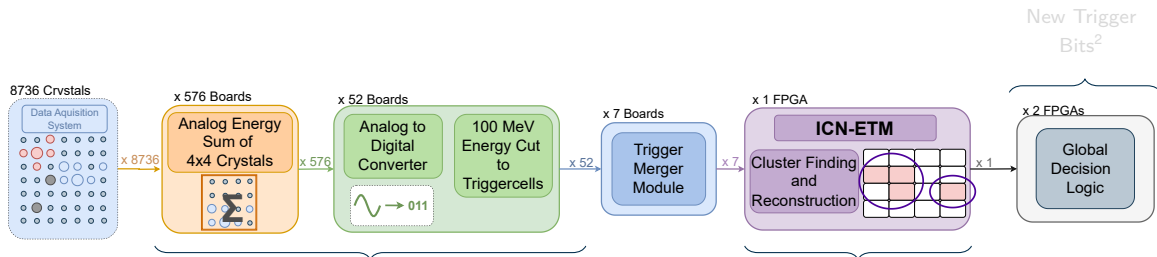
Improvement Possibilities after LS2



¹ GNNETM running in 2025c (paper in preparation, [Presentation at FastML](#))

² TauNN ([Presentation at FastML](#))

Improvement Possibilities after LS2



Increase granularity, improve feature extraction (ShaperDSP upgrade)
Change energy cut/crystal selection
(ML/Conventional algorithms for further filtering)

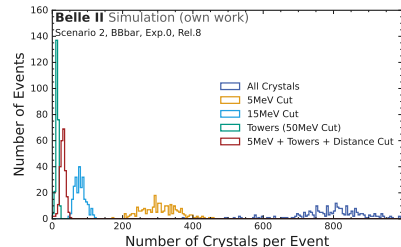
New clustering algorithms¹ with
increased latency budget

¹ GNNETM running in 2025c (paper in preparation, [Presentation at FastML](#))

² TauNN ([Presentation at FastML](#))

Improvements toward Upgrade - Input Possibilities

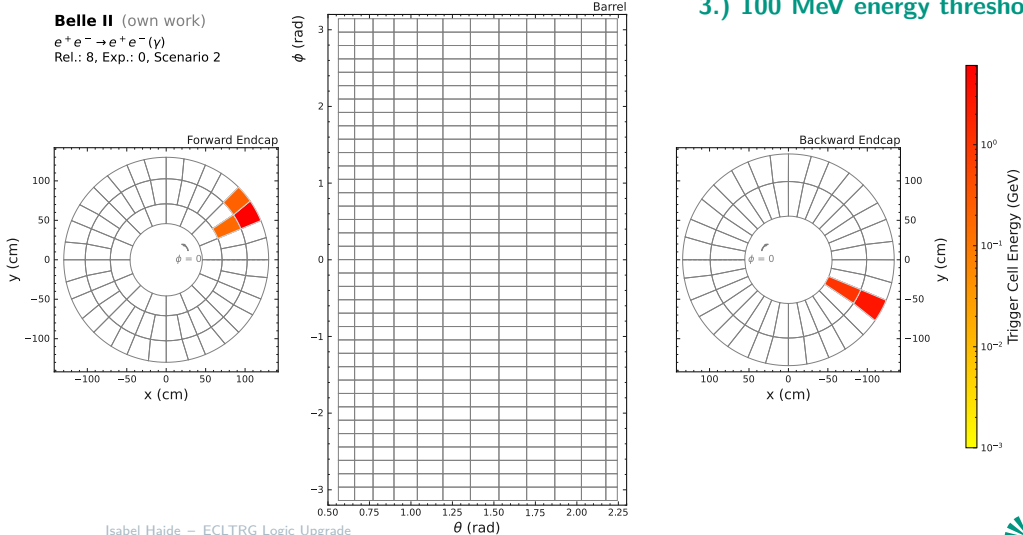
- ECLTRG input can be upgraded from 4x4 TCs to crystals with ShaperDSP upgrade. However, all active crystals (≈ 800 per event) as input is too much for even very modern hardware (Versal/AI Engines) with our current algorithms.
- We test two (so far purely theoretical) reduction possibilities to reduce the input to <128 :
 1. Flat energy cut for all crystals, which is similar to current handling of TCs.
 2. Selection of so-called trigger towers:
 - a. Find high energetic crystals, e.g. crystals above 50 MeV.
 - b. Select these crystals and a “tower” around it, all active crystals in a 5x5 area, if necessary with additional energy cut.
- One other possibility would be segmentation of ECL and processing on multiple FPGAs.



4x4 TC Event Display

- 1.) Timing window of 250 ns
- 2.) 4x4 energy sum
- 3.) 100 MeV energy threshold

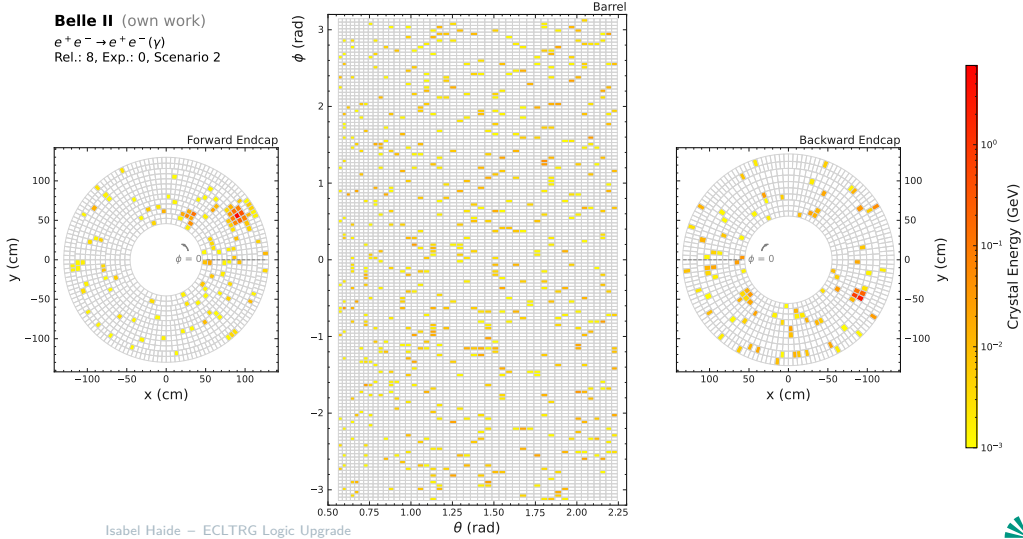
Belle II (own work)
 $e^+e^- \rightarrow e^+e^-(\gamma)$
 Rel.: 8, Exp.: 0, Scenario 2



1x1 TC Event Display

1.) Timing window of 250 ns

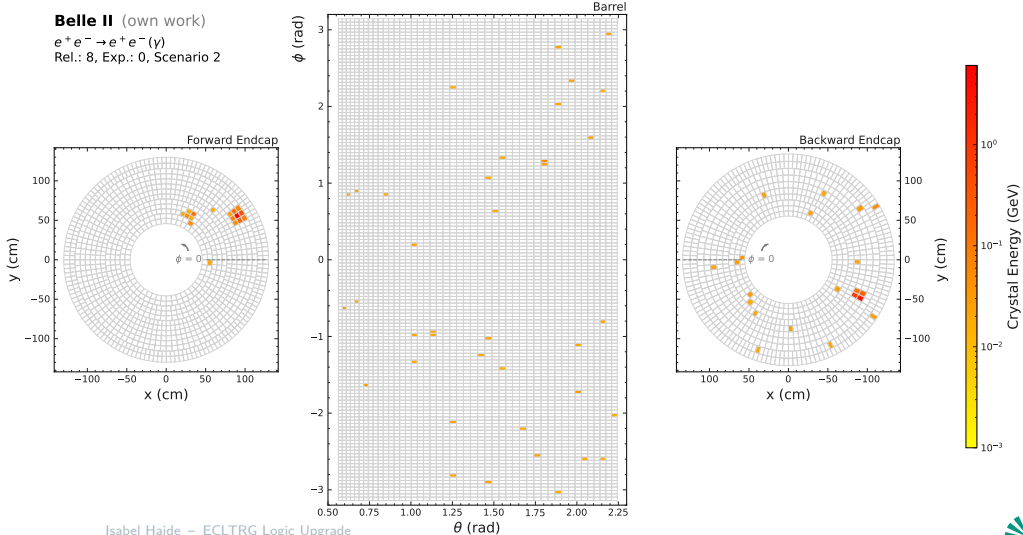
Belle II (own work)
 $e^+e^- \rightarrow e^+e^-(\gamma)$
 Rel.: 8, Exp.: 0, Scenario 2



1x1 TC - Flat Energy Cut

- 1.) Timing window of 250 ns
- 2.) 15 MeV energy threshold

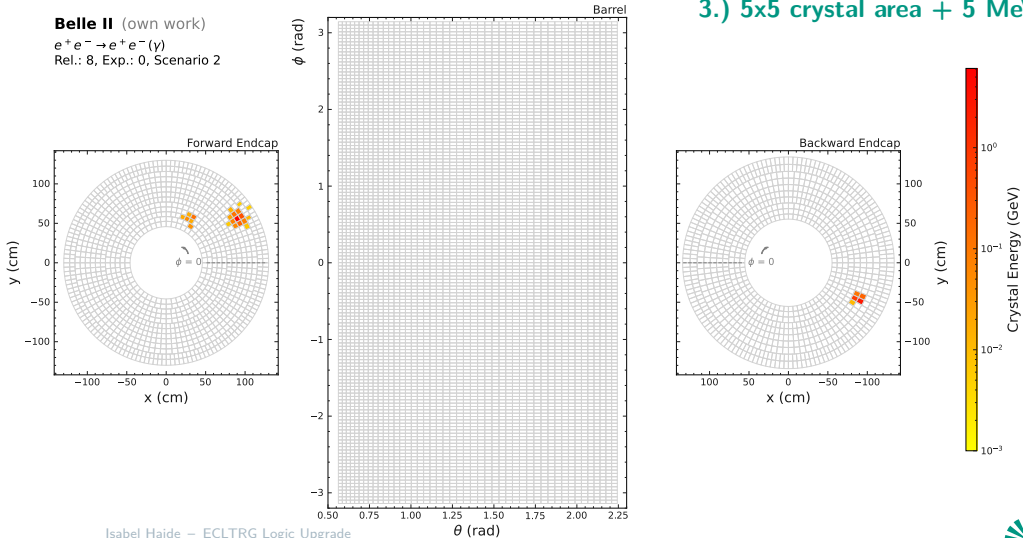
Belle II (own work)
 $e^+e^- \rightarrow e^+e^-(\gamma)$
 Rel.: 8, Exp.: 0, Scenario 2



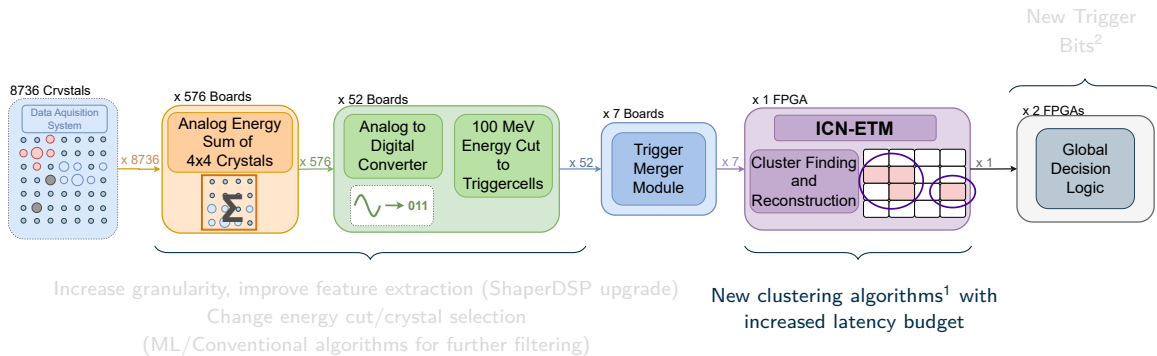
1x1 TC - Trigger Towers

- 1.) Timing window of 250 ns
- 2.) 50 MeV high energy cut
- 3.) 5x5 crystal area + 5 MeV cut

Belle II (own work)
 $e^+e^- \rightarrow e^+e^-(\gamma)$
 Rel.: 8, Exp.: 0, Scenario 2



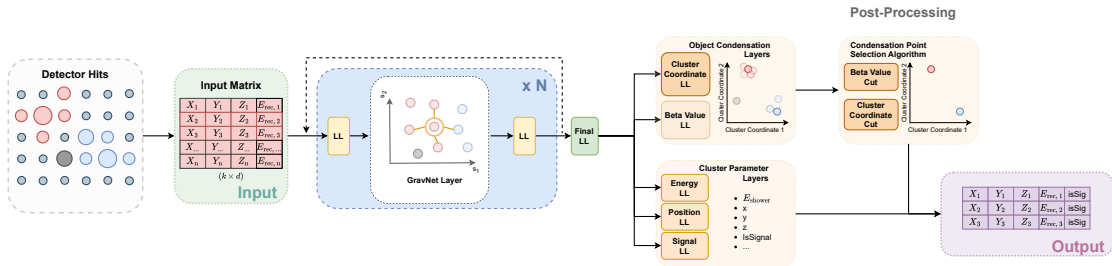
Improvement Possibilities after LS2



¹ GNNETM running in 2025c (paper in preparation, [Presentation at FastML](#))

² TauNN ([Presentation at FastML](#))

Improvement Possibilities - Clustering Algorithm



- We need a working algorithm to process and cluster 1x1 crystals within the restrictions of the post-LS2 trigger system (≈ 8 MHz throughput, $10 \mu s$ latency).
- GravNet architecture has proven to work well for clustering algorithms in the ECL for both offline¹ and online² applications.
- We use the same architecture as for the GNN-ETM model, but applying it to single crystals with flat energy cut/trigger tower selection.

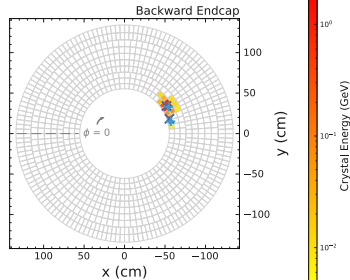
¹ Photon Reconstruction in the Belle II Calorimeter Using Graph Neural Networks, F. Wemmer et al. ([Paper](#))

² GNNETM (paper in preparation, [Presentation at FastML](#))

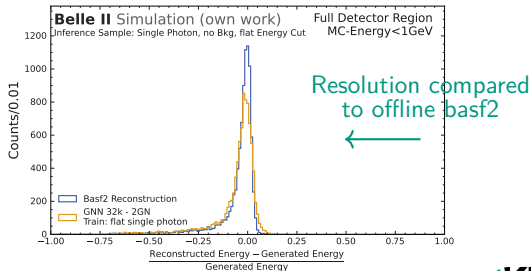
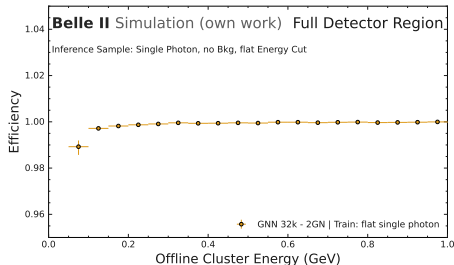
Clustering Algorithm - Proof of Concept

- First trained models on very simplified datasets (no/low beam background, particle gun photons) show promising performance.
- Model is not optimized for hardware implementation, but has low number of parameters (<5000).
- We are moving towards more realistic scenarios and testing the different input reduction possibilities.

+ GNN Prediction
x Offline Clusters

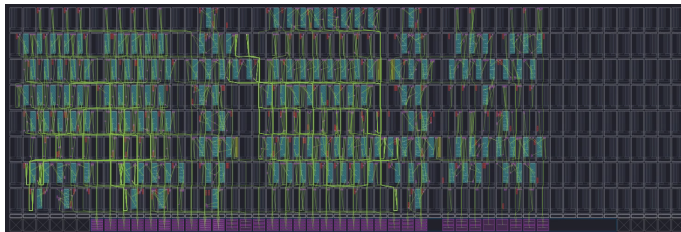


Single Photons, no beam background



Clustering Algorithm - Implementation

- We test the GNN network with a maximum of 128 inputs on an AMD Versal VCK190 Board, featuring a XCVK1902 with 400 AI Engines. The high number of inputs is very challenging even for this hardware, we use a hybrid implementation with both AI Engines and the programmable logic.
- The first full, successful implementation reaches a latency of $\approx 10 \mu\text{s}$.
- The programmable logic is running at $f_{\text{sys}} = 256 \text{ MHz}$, the total throughput is at $\approx 1 \text{ MHz}$. A new tiling concept for the AIEs could increase the throughput up to 4 MHz by reducing memory stalls.



	HLS Part	AIE Part
LUT	55.74 %	0 %
FF	25.37 %	0 %
DSP	17.07 %	0 %
BRAM	67.43 %	0 %
AIE tiles for calculations	0 %	40 %
Total used AIE tiles	0 %	66.75 %

Summary and Outlook

- New readout possibilities due to ShaperDSP upgrade increase granularity and allow use of shower shape information, but necessitate the design of new algorithms.
- Additionally, higher luminosity requires new ECL clustering algorithms to avoid prescaling of trigger bits.
- Readout of all crystals not possible to process with our algorithms on our current hardware → Reduction of input size with flat energy cut or trigger tower implementation
 - Which is possible on the hardware? Could the trigger towers be implemented?
- We show a first proof of concept ML algorithm for a one-step clustering, which on very simplified events can provide a good efficiency and resolution.
- We implement the model design on an AMD Versal VCK190 Board with an XCVC1902, reaching a latency of 10 μ s. To reach the necessary throughput for the Belle II trigger post-LS2, further optimizations are necessary.