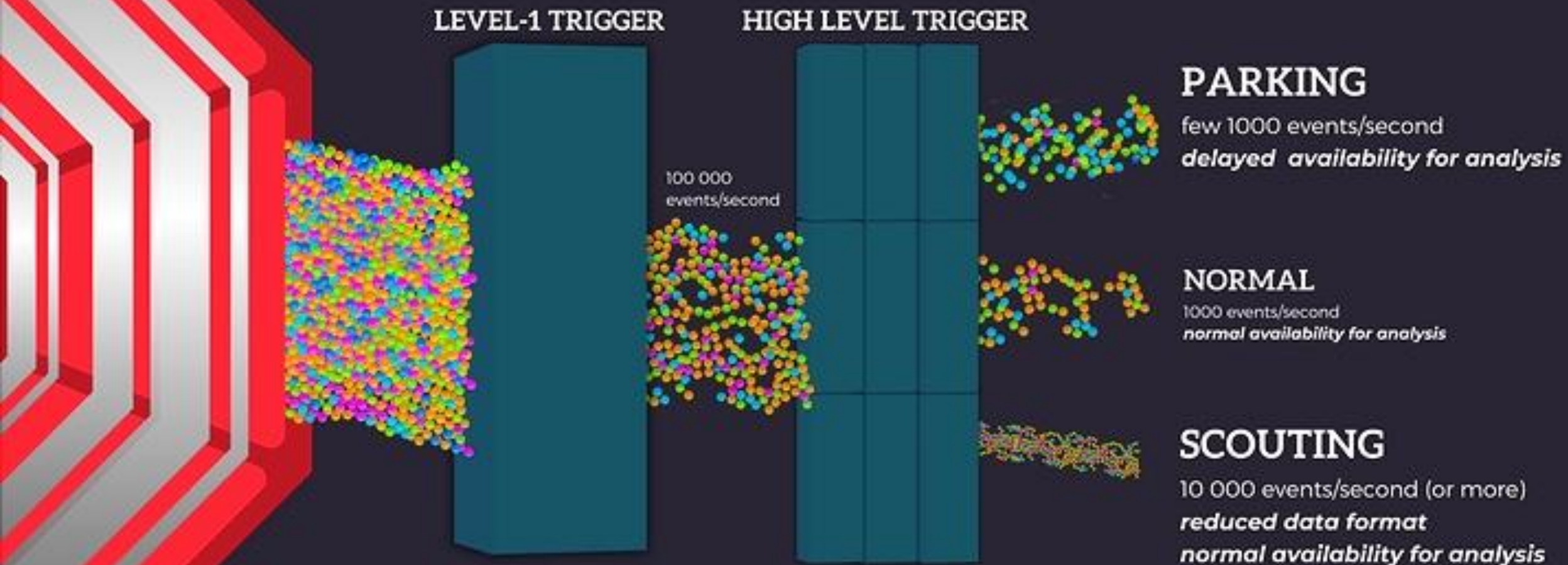


Exploring machine learning to reconstruct electron and photon energies in the CMS scouting data

Bartosz Dlubak

Supervisor: A.R.Sahasransu

Objectives and CMS scouting data



CMS DETECTOR

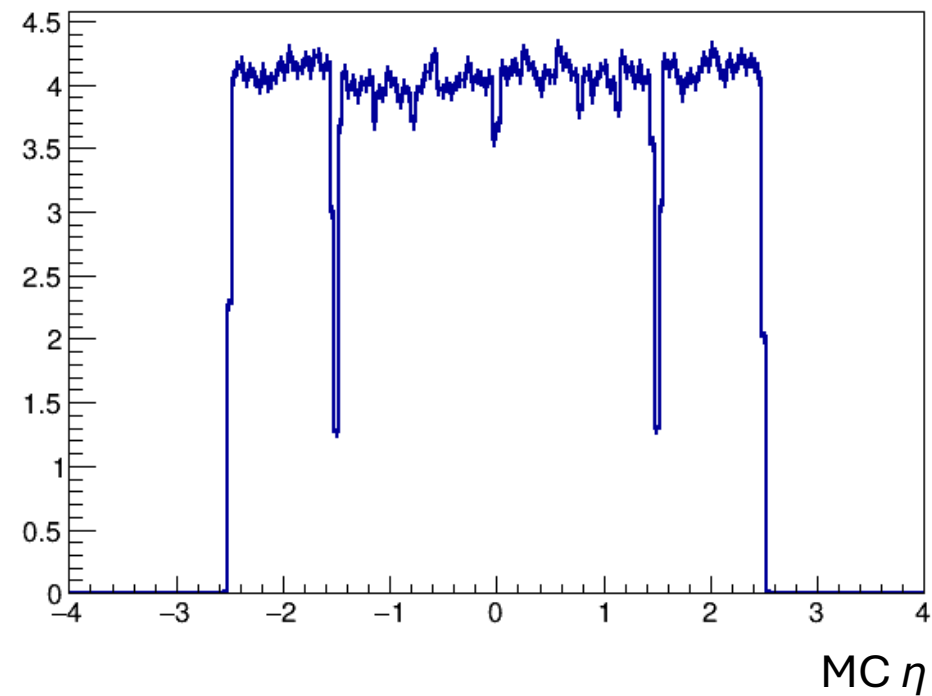
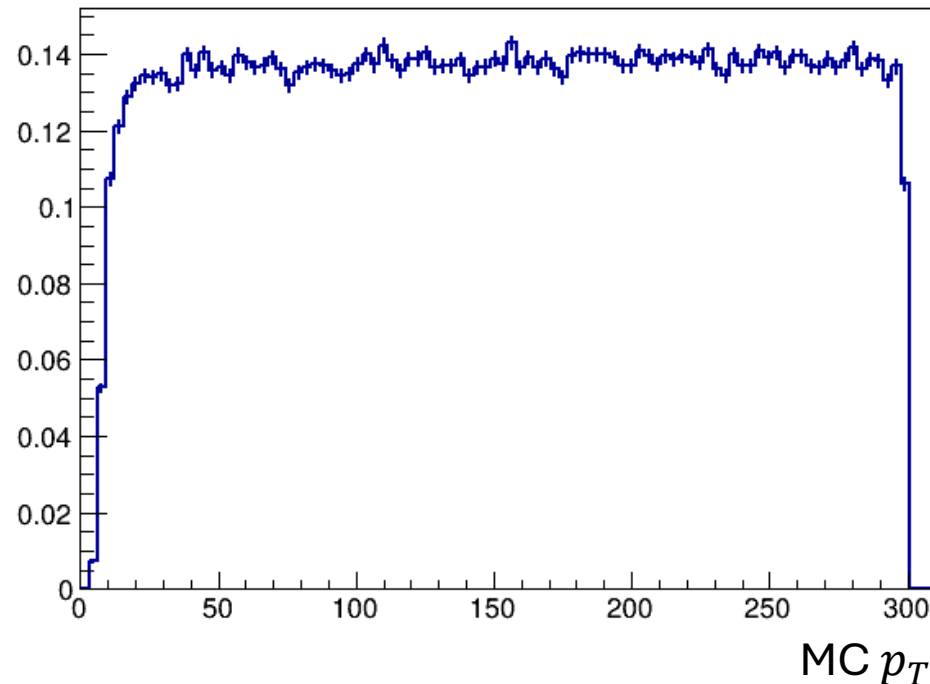
records 40 000 000 times/second

Aim: Reduce uncertainty in e/ γ energies in the CMS scouting data.

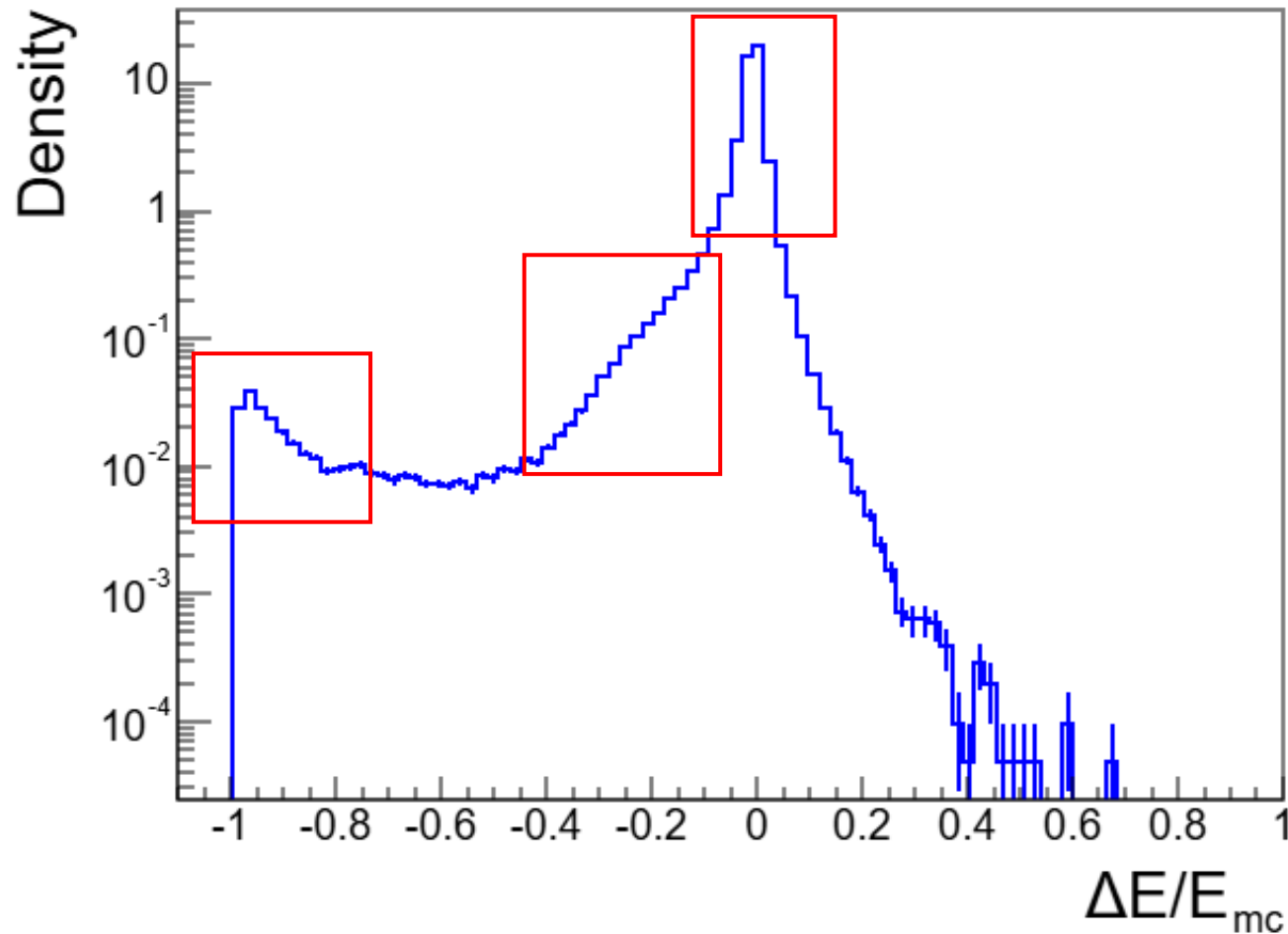
Credit: CMS Collaboration, F. Blekman

Simulated data

- Uniform distribution over MC electron (p_T , η , ϕ).
- Dips in MC η correspond to the barrel – endcap transition region.



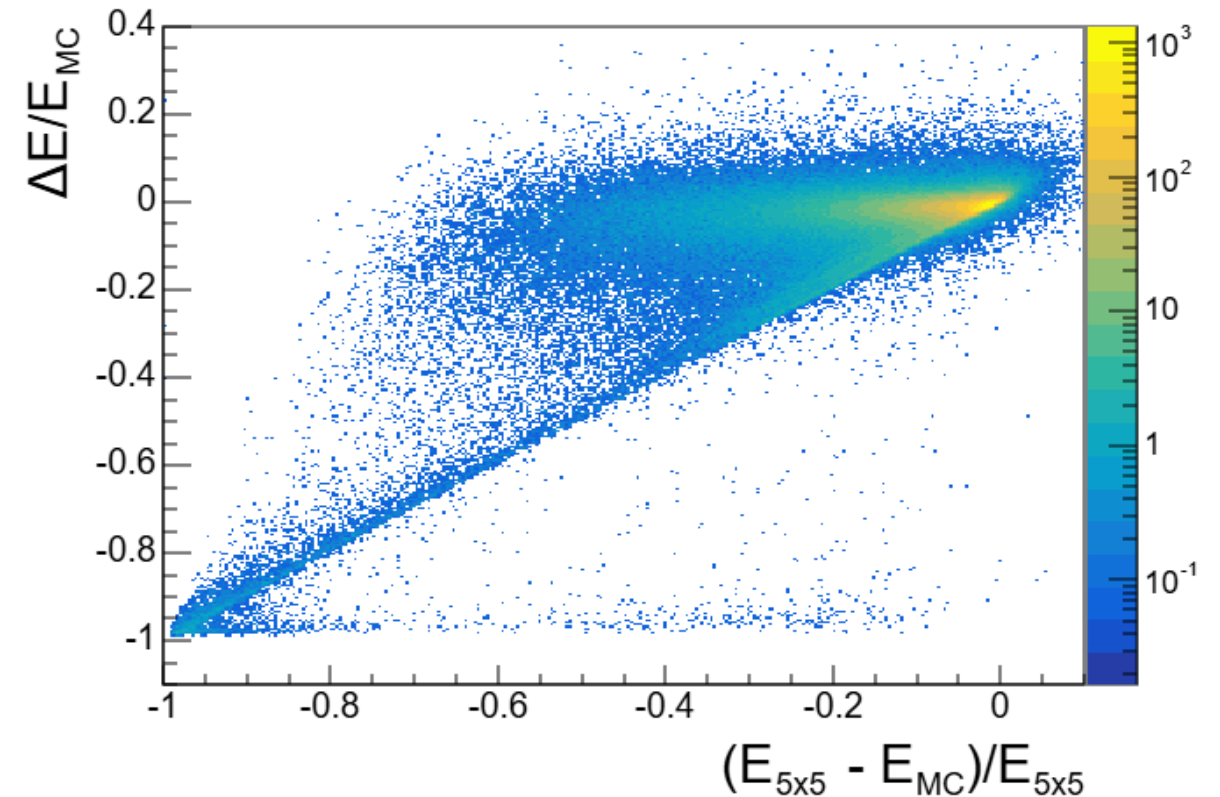
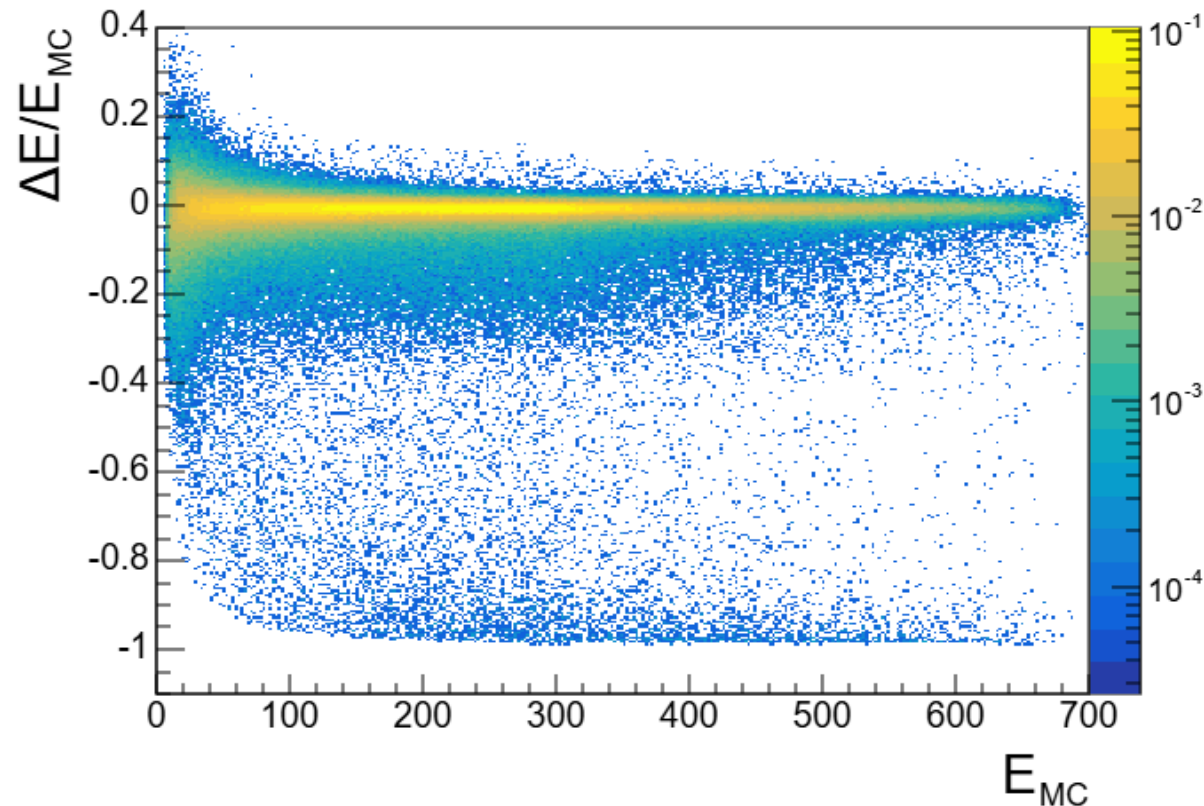
Uncertainty in scouting electron energy



- E_{reco} : calibrated ECAL energy
- $\Delta E/E_{MC} = ((E_{reco} - E_{MC})/E_{MC})$
- Peak observed ~ 0.0 ,
- Shoulder ~ -0.2 and
- Peak at -1.0

Correlations between features and uncertainty

- Shoulder contribution has a linear dependence on shower containment in a 5x5 crystal block

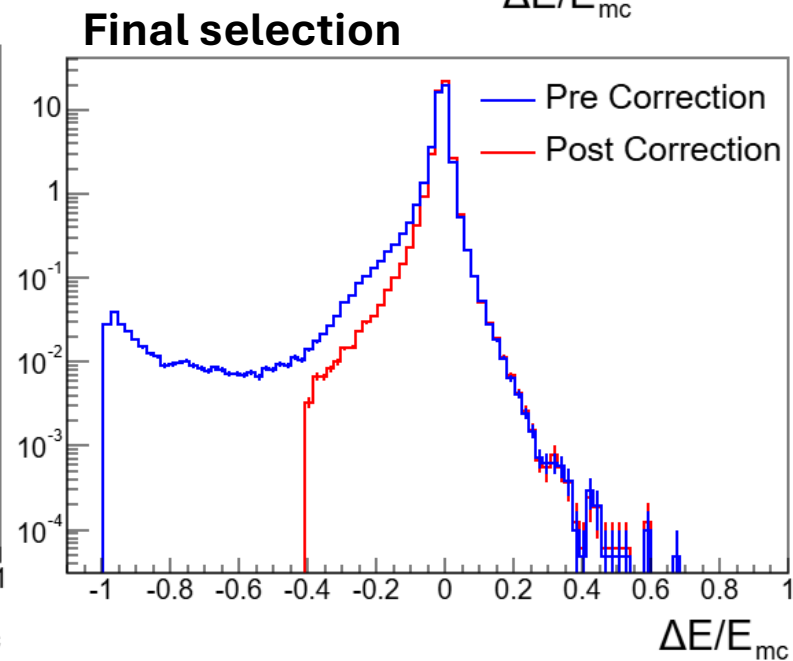
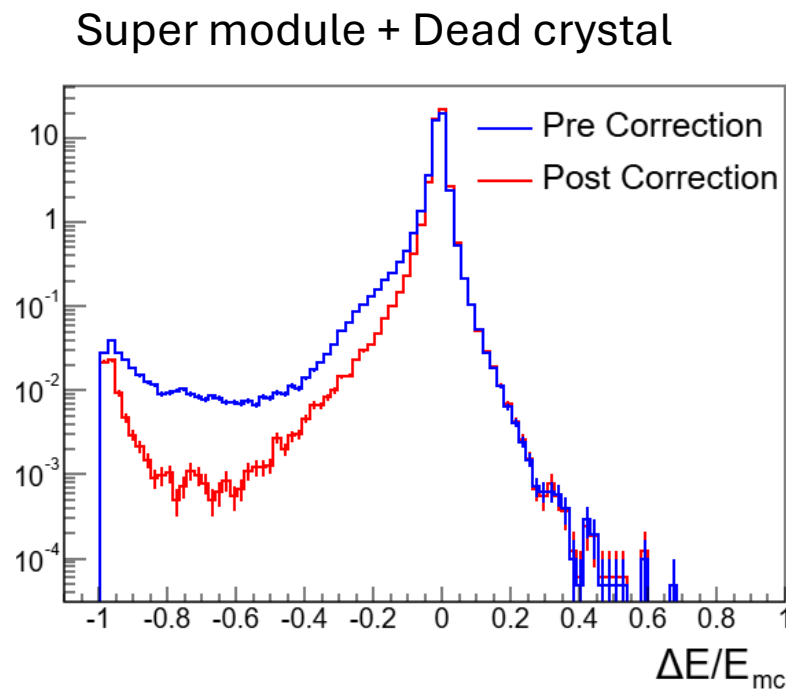
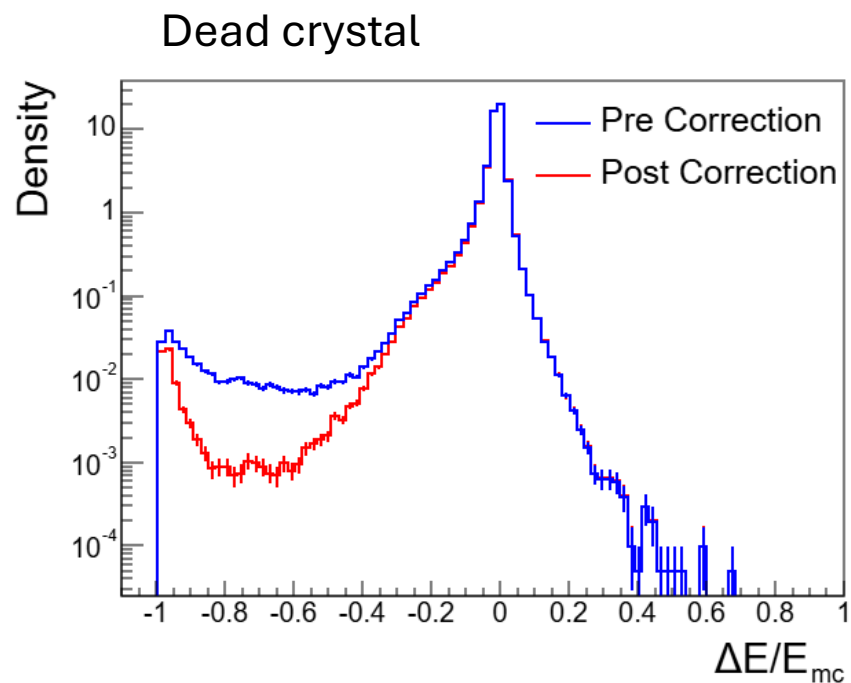
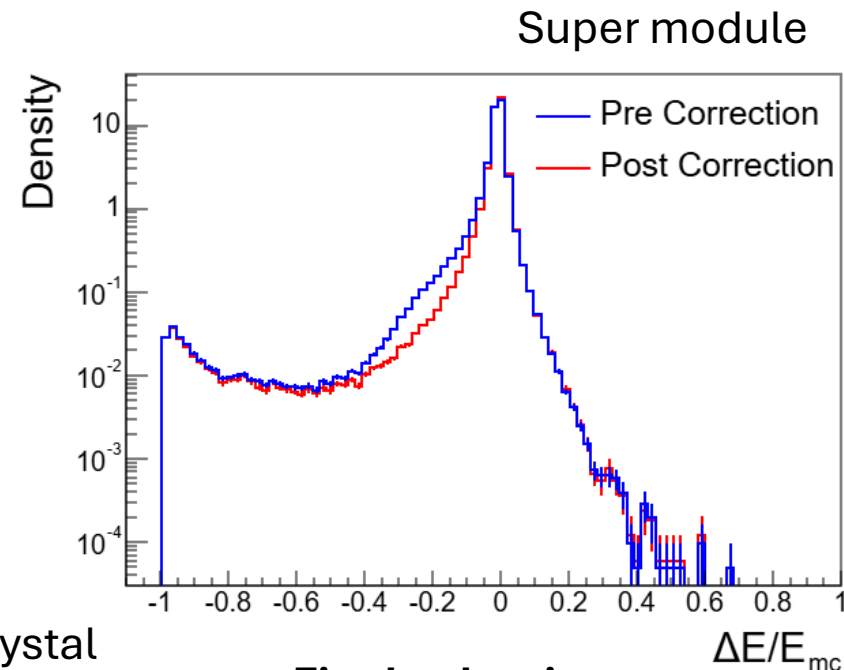


Identifying good e (Barrel)

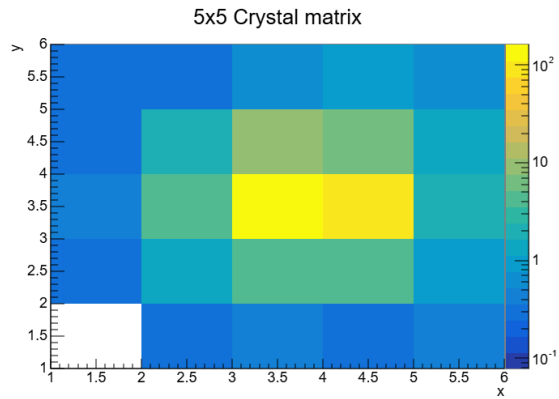
Two known noisy effects due to:

- Super modules and
- Dead crystals

Can be identified in the (η, ϕ) map.



ML Models: BDT and NN



Shower
shapes

(p_T, η, ϕ)

Tracker
data

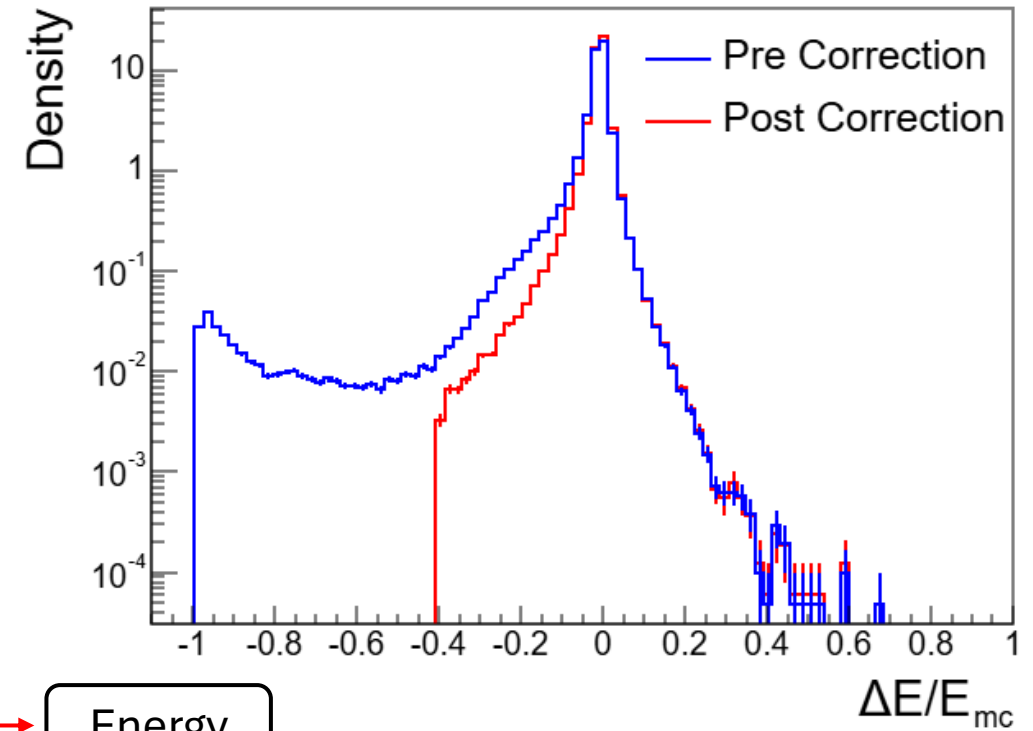
Trying two different
approaches:

NN

BDT

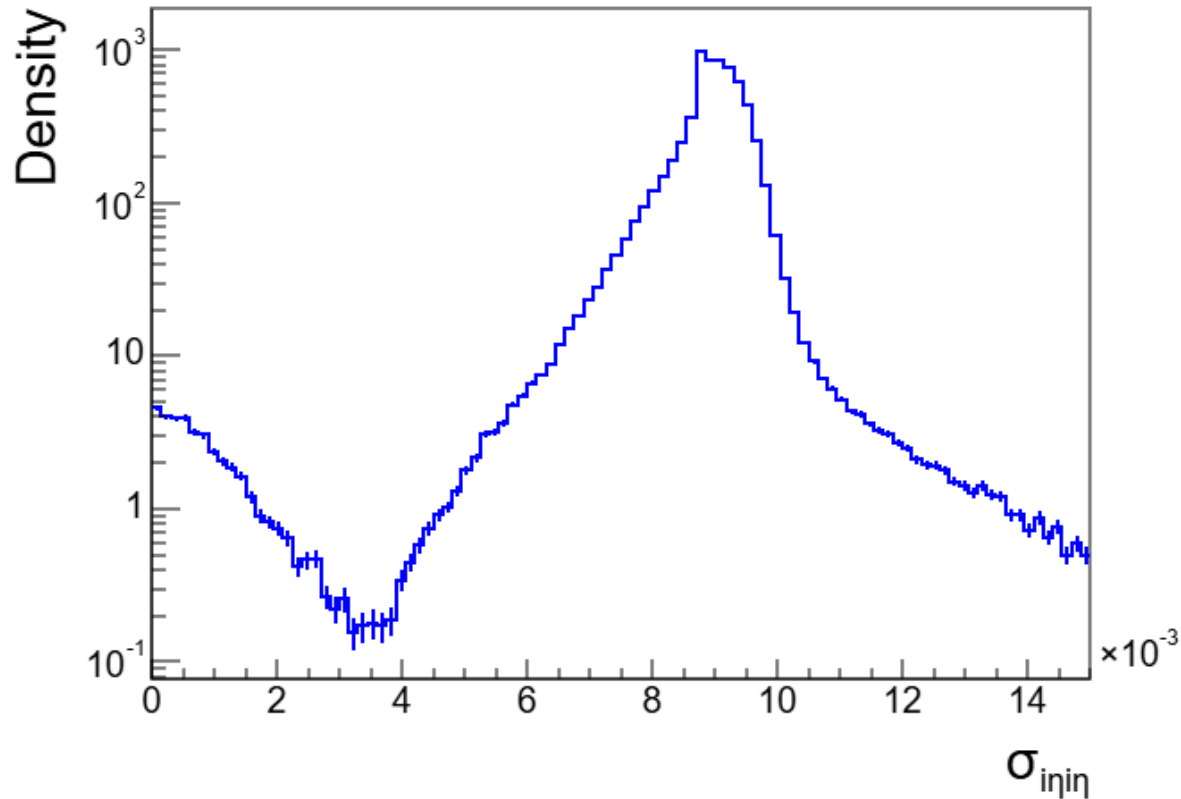
Energy

Correction
factor



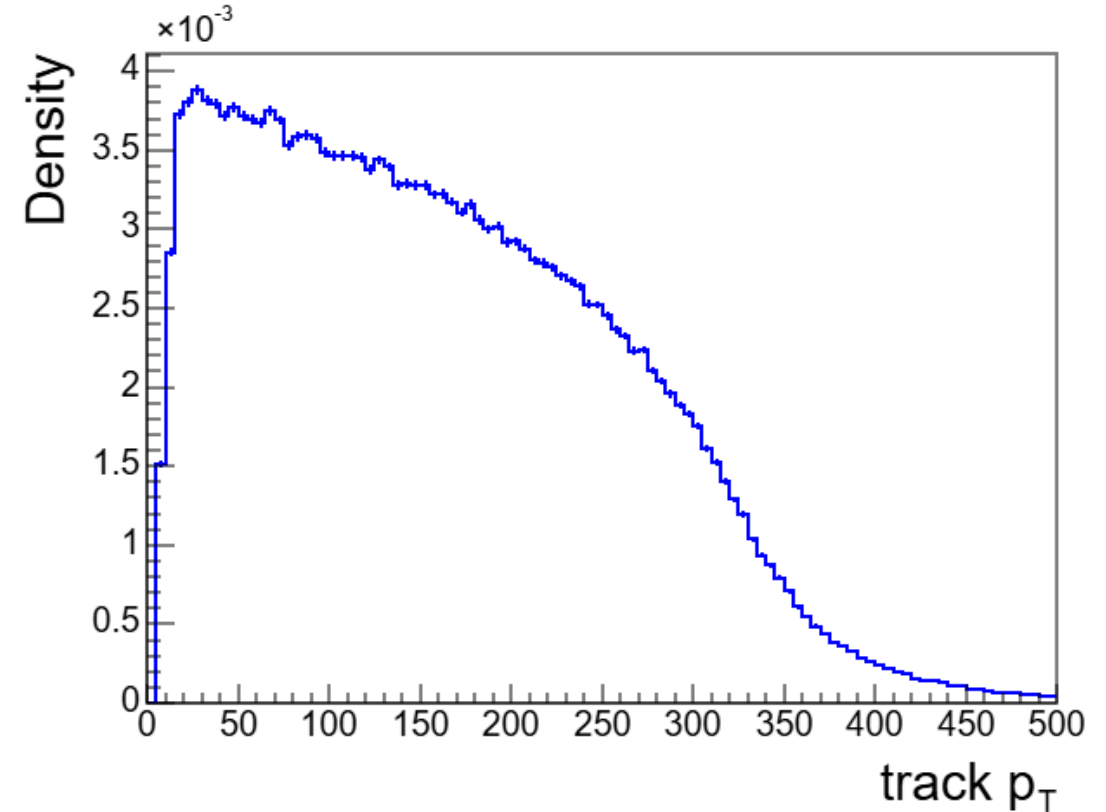
Use multivariate
analysis models to
reduce the uncertainty
on the energy
reconstruction for e/γ .

Input features



$$\sigma_{i\eta i\eta} = \sqrt{\frac{\sum_i^{5 \times 5} w_i (\eta_i - \bar{\eta}_{5 \times 5})^2}{\sum_i^{5 \times 5} w_i}}$$

$$w_i = \max(0, 4.7 + \ln(E_i/E_{5 \times 5}))$$

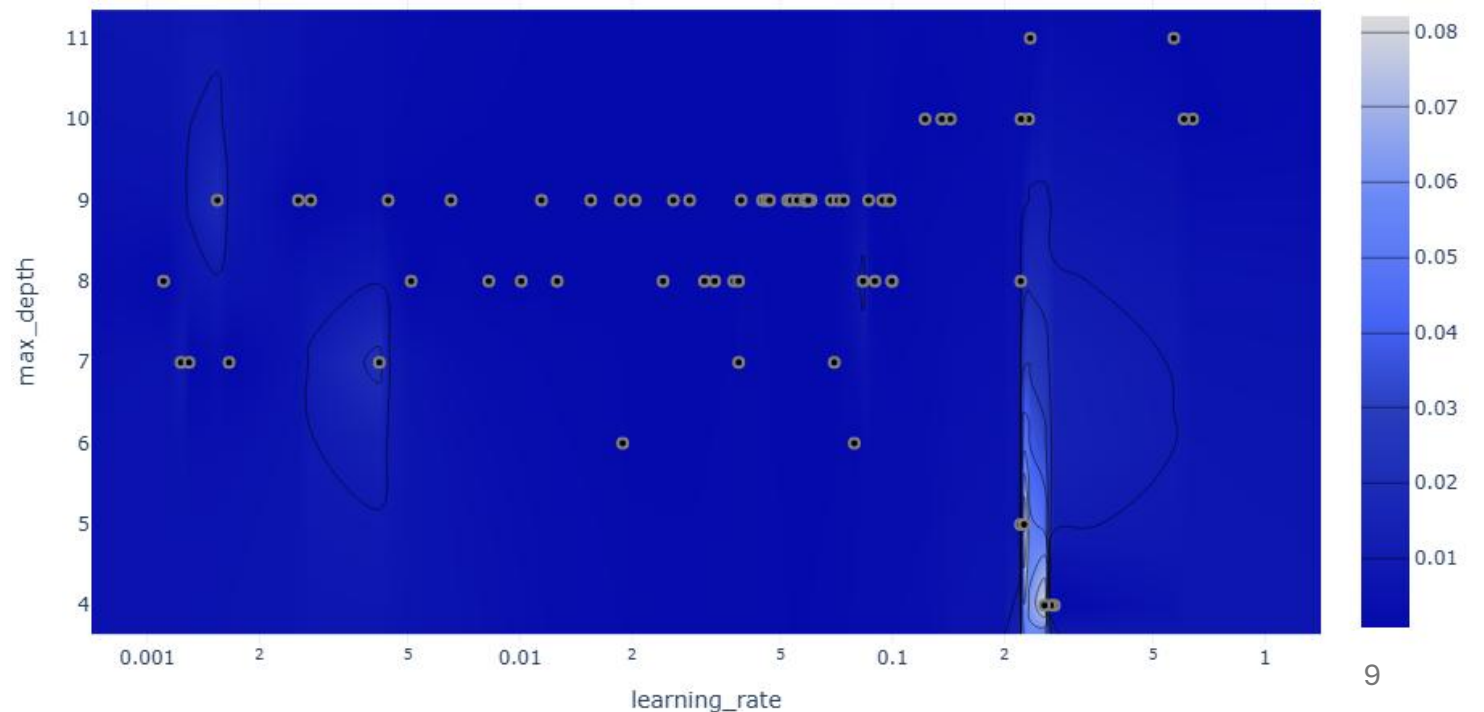
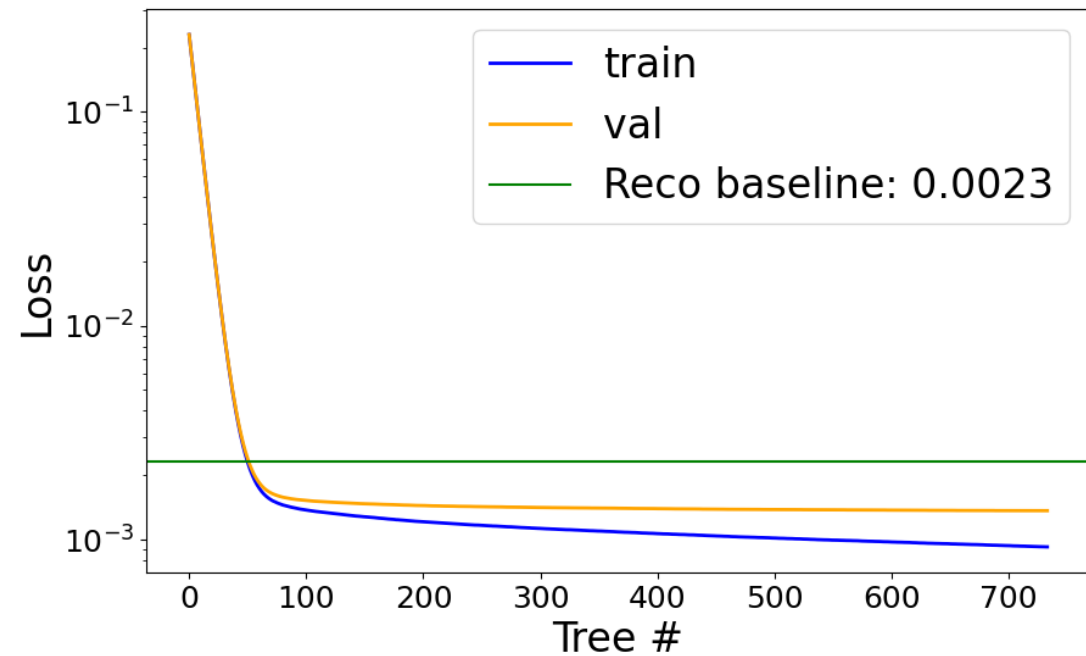


Two distinct category of features;

- ECAL only for γ
- Additional track parameters for e

BDT Training

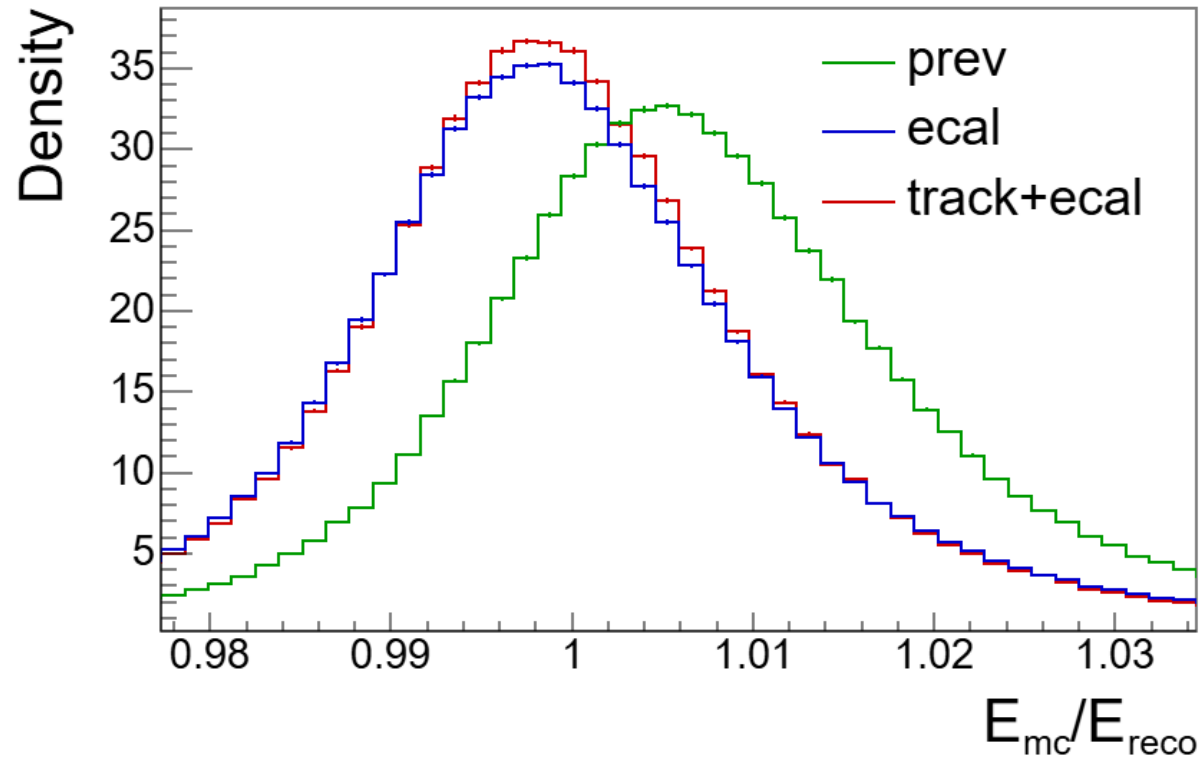
- Targets:
 - $y_{true} = E_{MC}/E_{reco}$
- Loss:
 - $= ((E_{pred} - E_{MC})/E_{MC})^2$
 - $= (\frac{y_{pred}}{y_{true}} - 1)^2$
- Used Optuna for hyperparameter optimisation.
 - <https://optuna.org/>



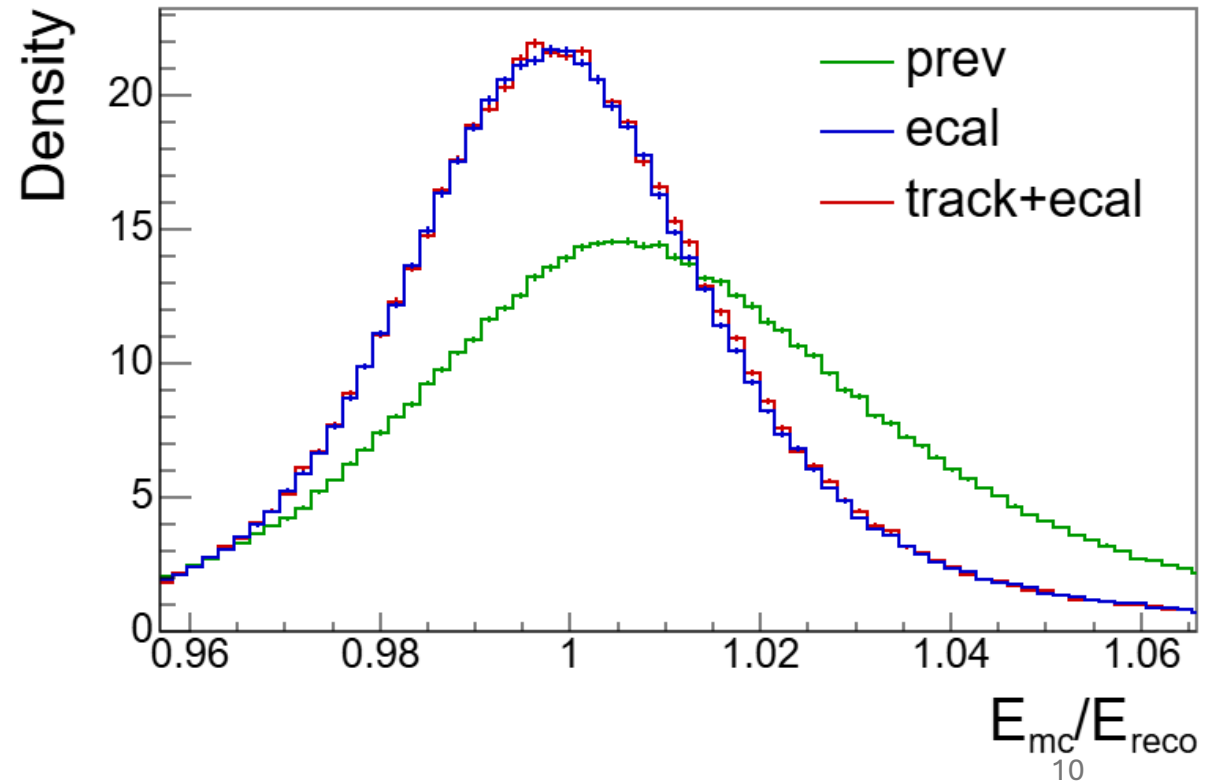
BDT Results

- Peak shifted closer to one and lower spread in both ECAL barrel and endcap.

Barrel (EB)

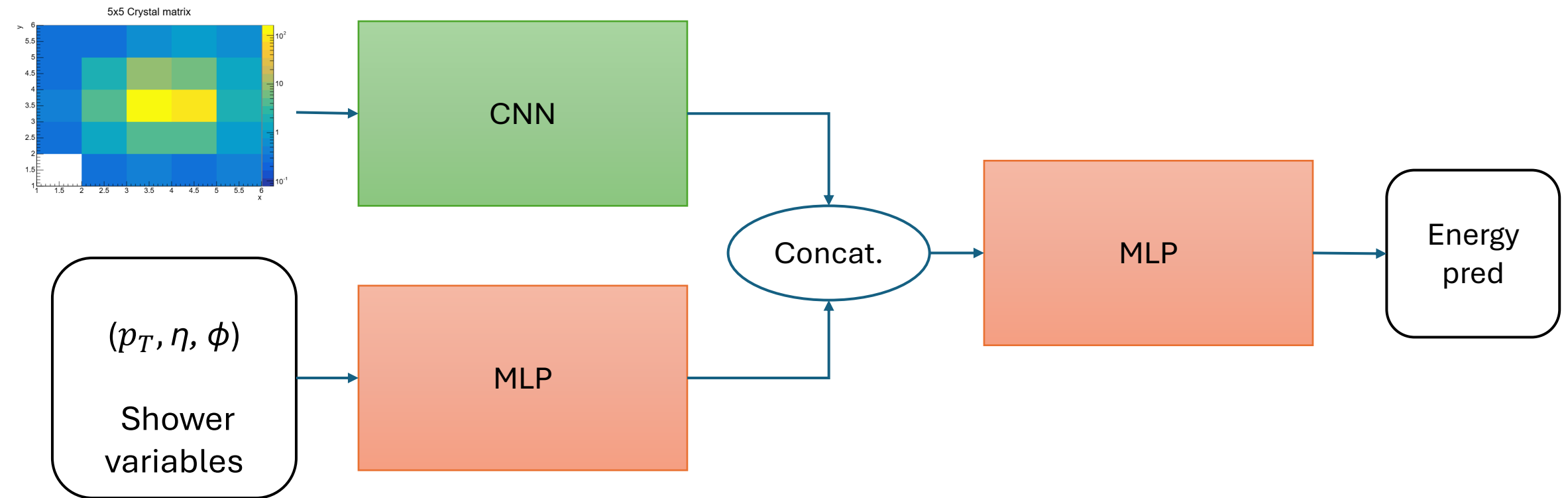


Endcaps (EE)



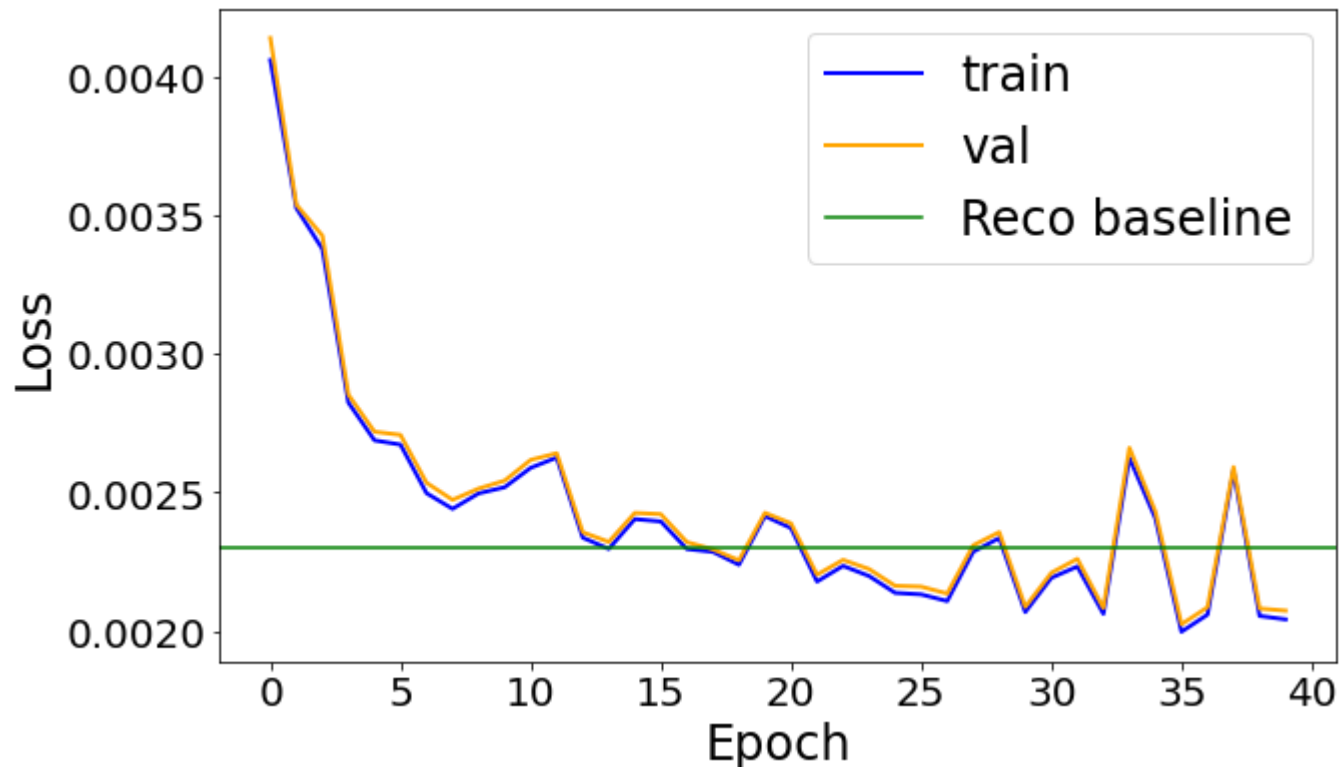
Deep Learning: CNN based architecture

- $Loss := ((E_{pred} - E_{MC}) / E_{MC})^2$



NN Training

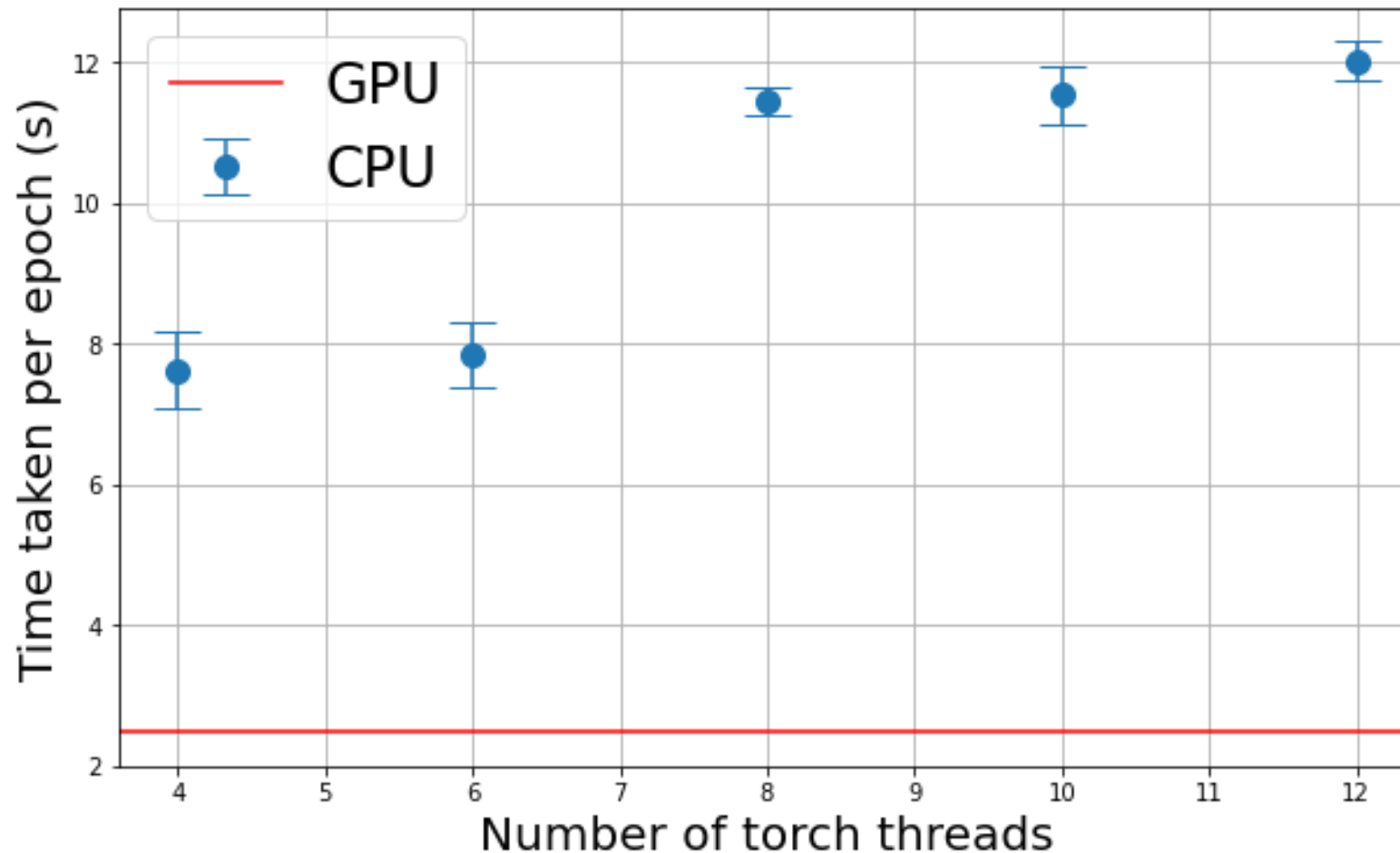
- Validation and train loss curve are very similar
 - =>Train and validation come from a similar distribution (large enough datasets)



- $Loss := ((E_{pred} - E_{MC}) / E_{MC})^2$

Acceleration with GPUs for NNs

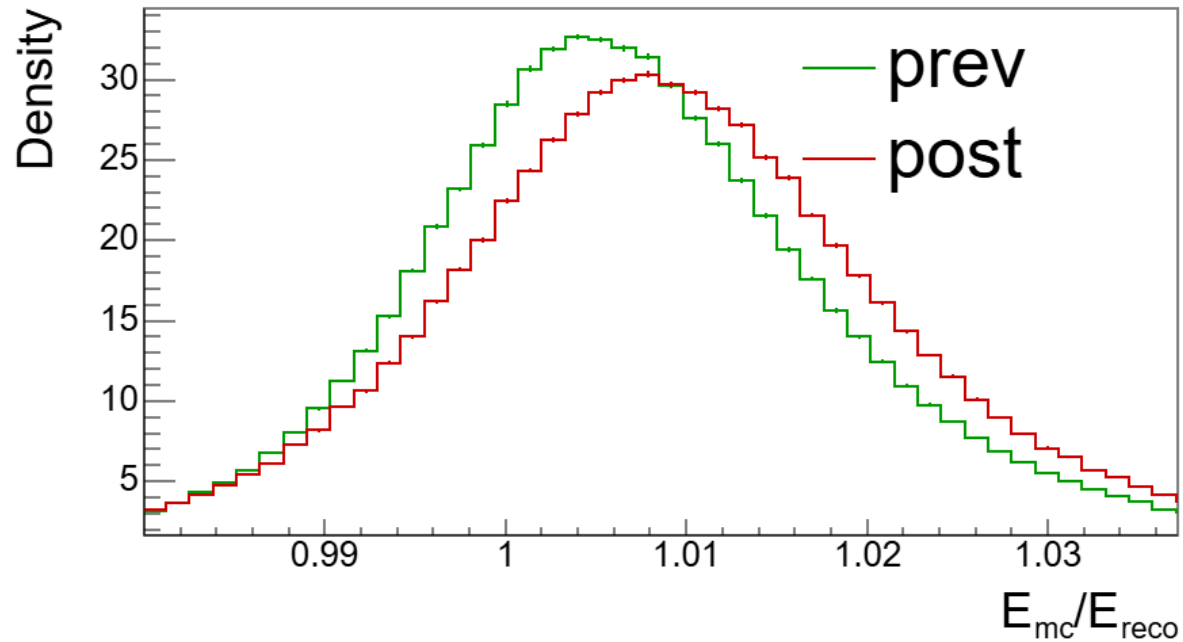
- GPU 4-6x faster than all CPU setting



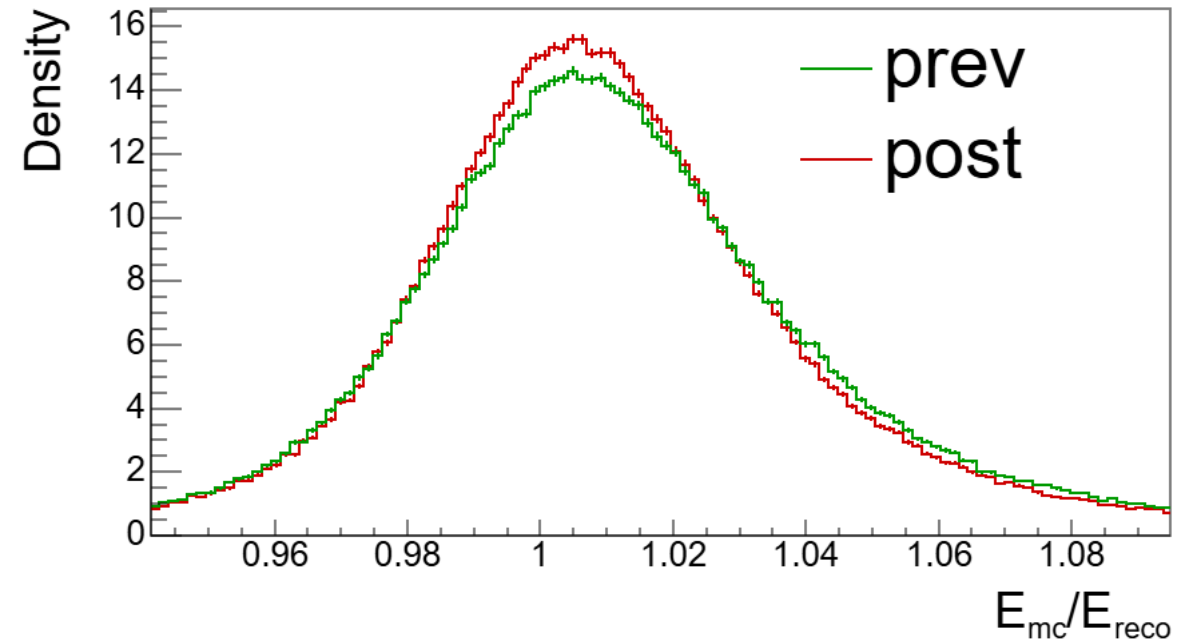
Results of NN

- CNN needs to be further optimised

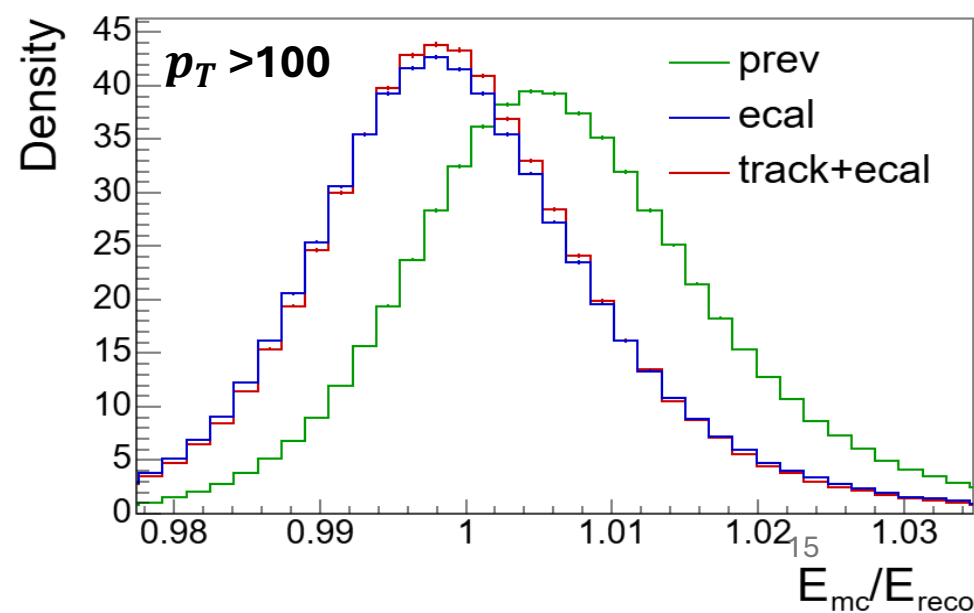
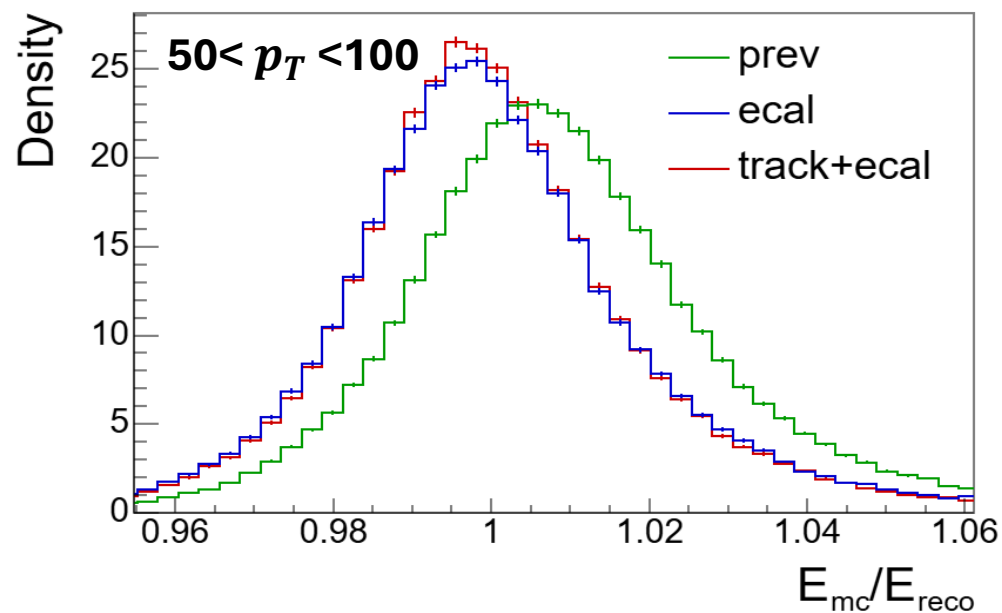
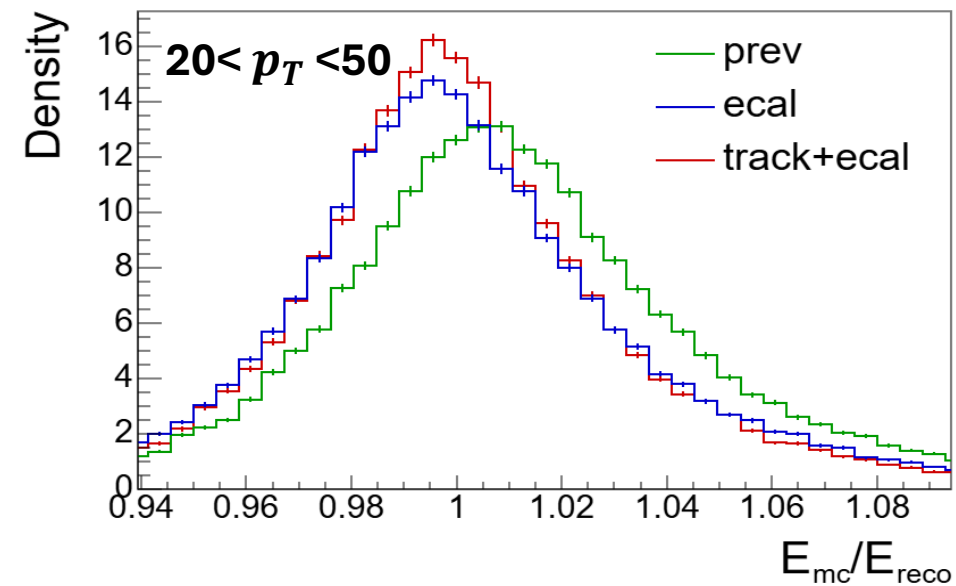
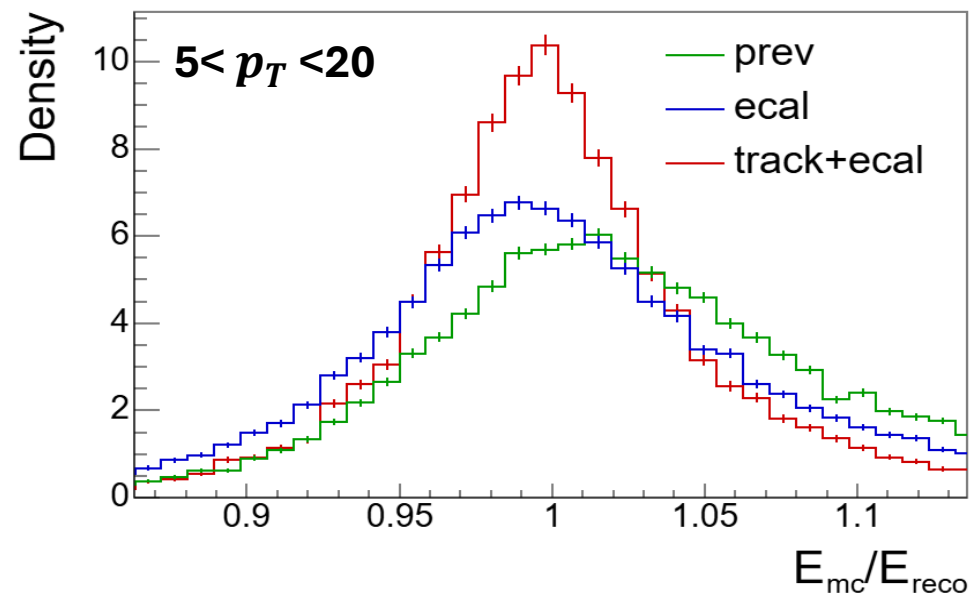
Barrel (EB)



Endcaps (EB)



Lowest uncertainty results (BDT)



Conclusions and outlook

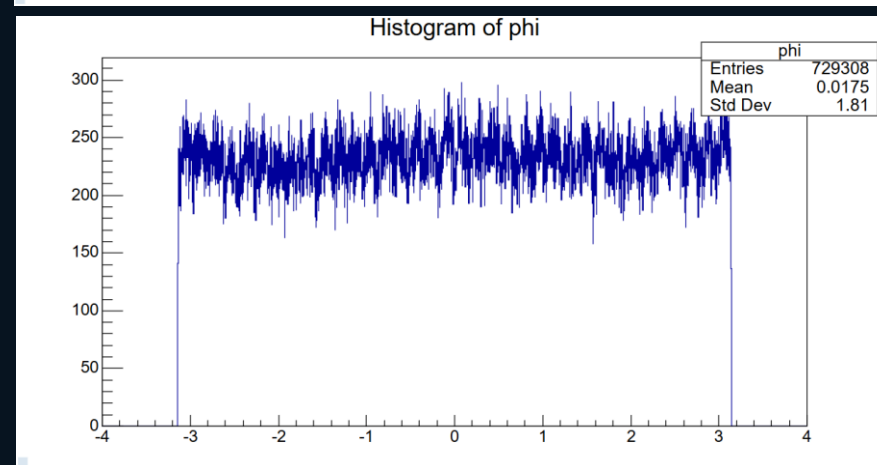
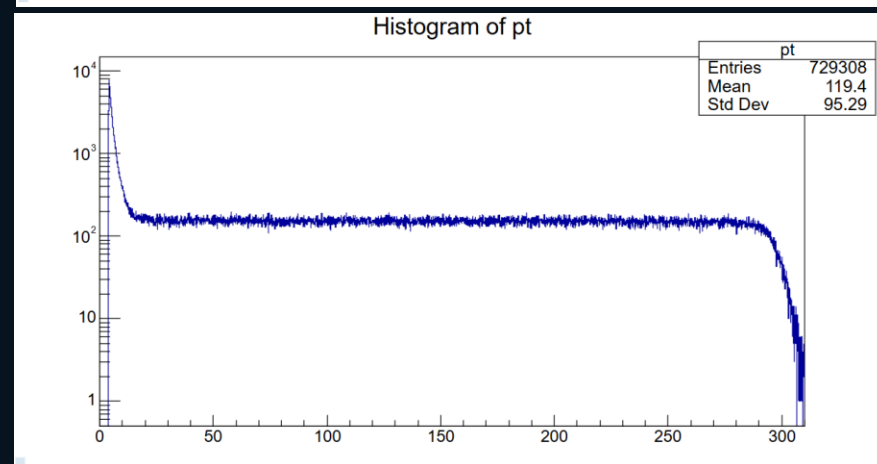
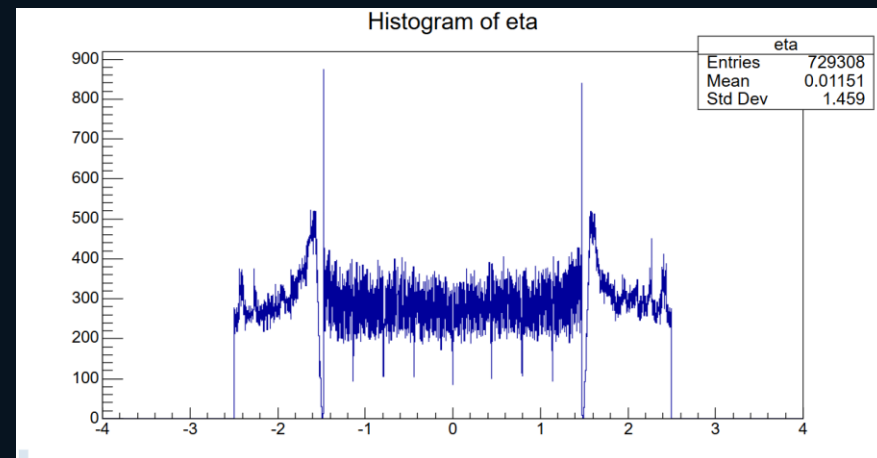
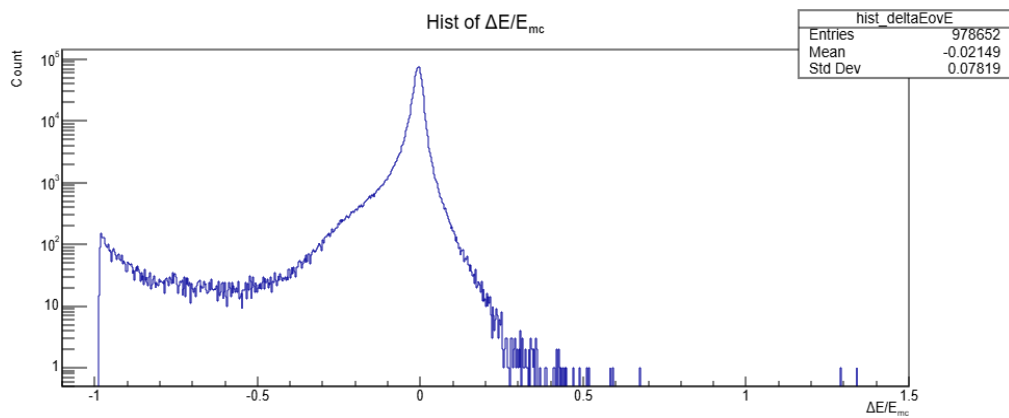
- Analysed MC e/γ sample and used ML to reduce the uncertainty on scouting data.
- Largest reduction in uncertainty observed for electrons in $5 < p_T < 20$ GeV.
- Consistent reduction in energy uncertainty observed for electrons and photons through all energies.
- Electron/photons incident on super module gaps need to be addressed separately.
- CNN approach needs to be further optimised
 - Shower shape variables, e.g. $sieie$, seemed to outperform the CNN.

BACKUPS SLIDES

Target data

- Monte Carlo simulations produce generated particles
- Each particle consists of three kinematic variables:
 - ϕ : Azimuthal angle $[-\pi, \pi]$
 - η : Pseudo rapidity
 - p_T : Transverse momentum
- Can combine p_T and η to get energy:
 - $E_{MC} = \sqrt{(p_T \cosh \eta)^2 + m^2}$, the Monte-Carlo simulated energy

Preliminary predictions

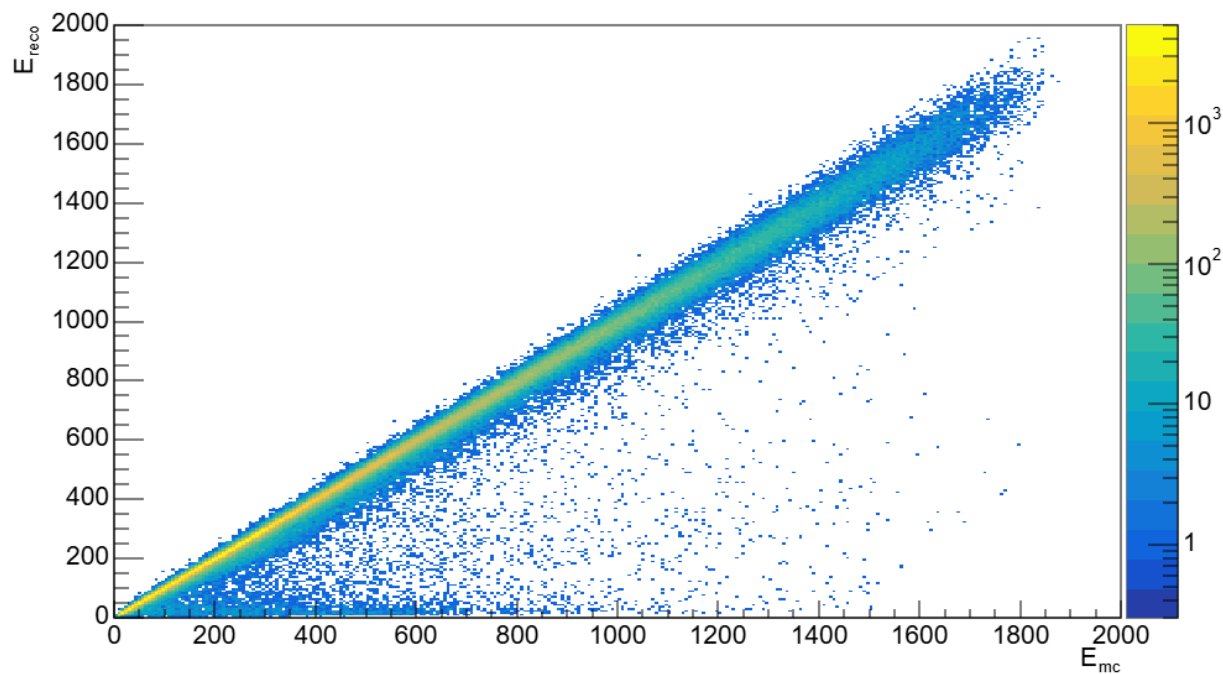


Investigating large error regions

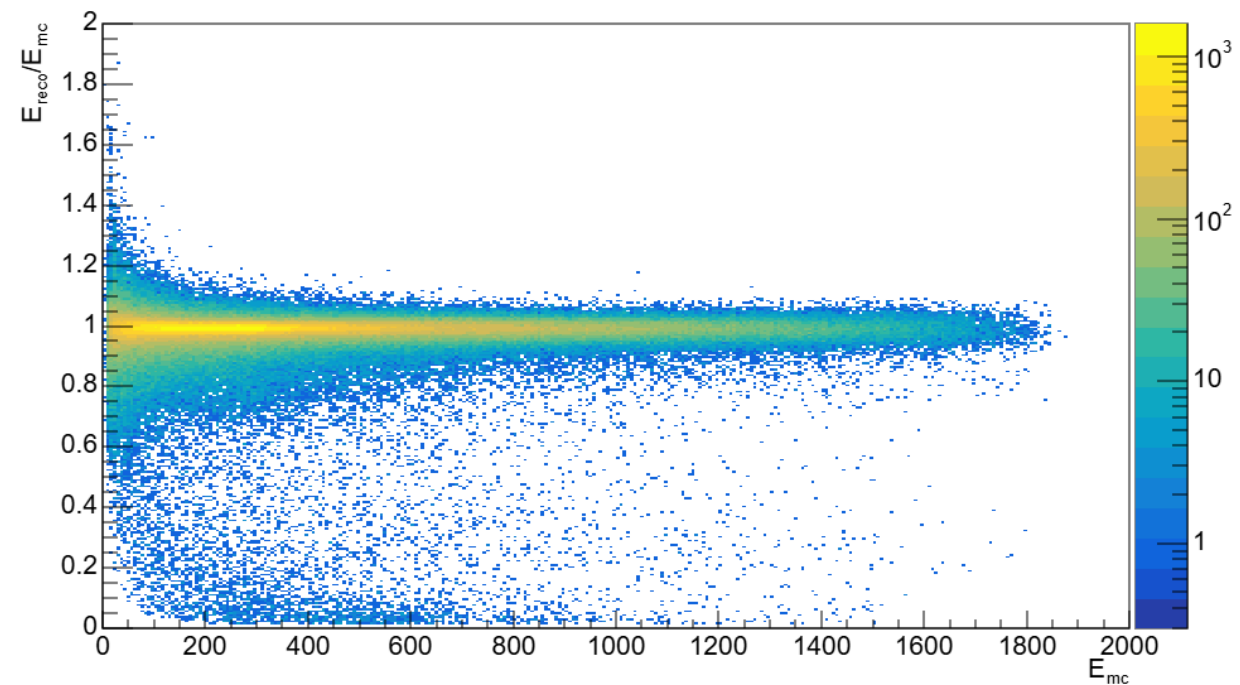
- Instead of decaying, the left tail of the $\Delta E/E_{MC}$ distribution grows significantly.
- Suggests $E_{reco} = 0$, as $\frac{0 - E_{MC}}{E_{MC}} = -1$.
- Critical to investigate the origin of this effect, as it may introduce systematic biases and degrade ML model performance.

E_{reco} vs E_{MC}

Density map of E_{reco} vs E_{mc}



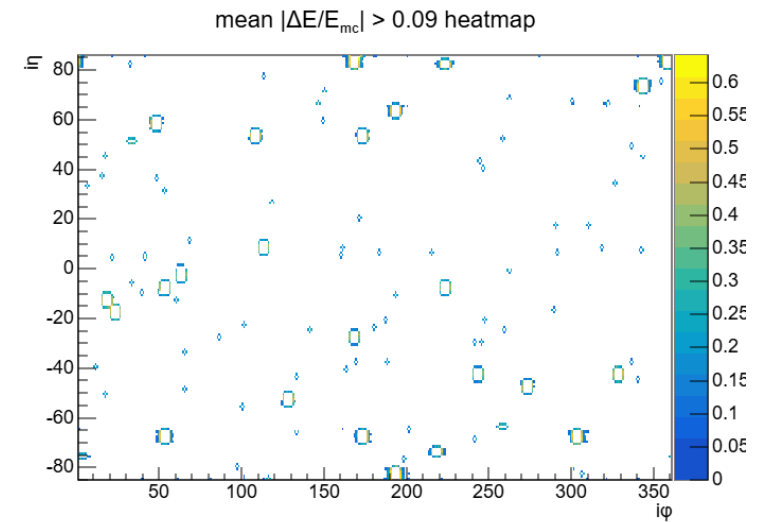
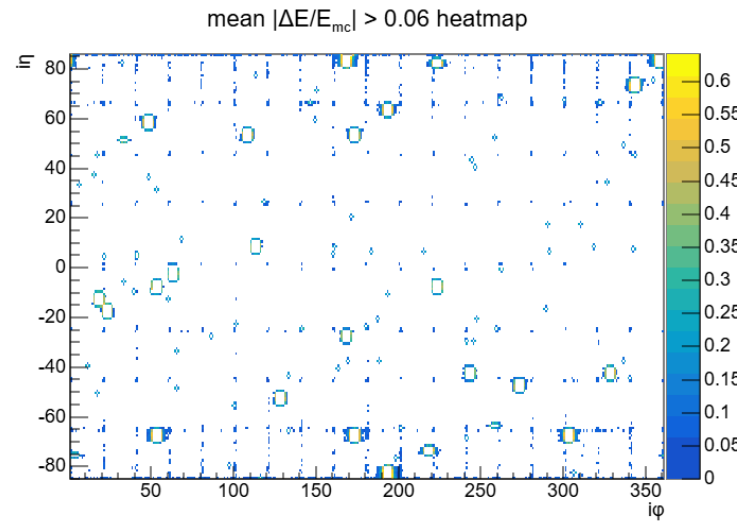
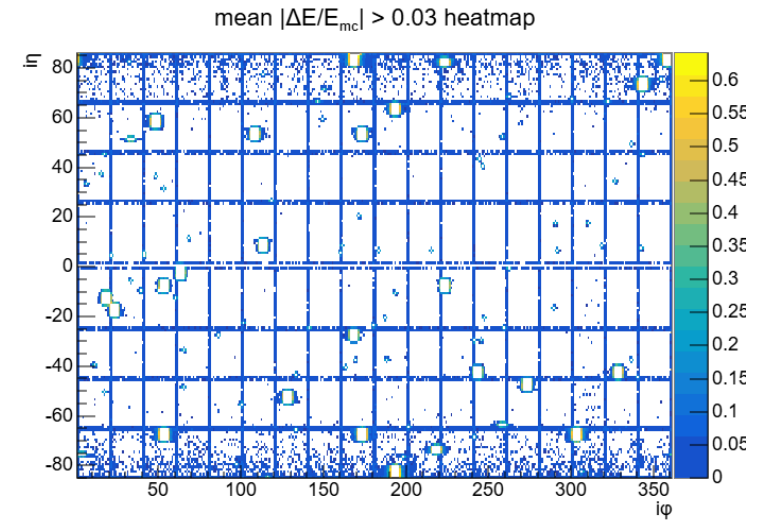
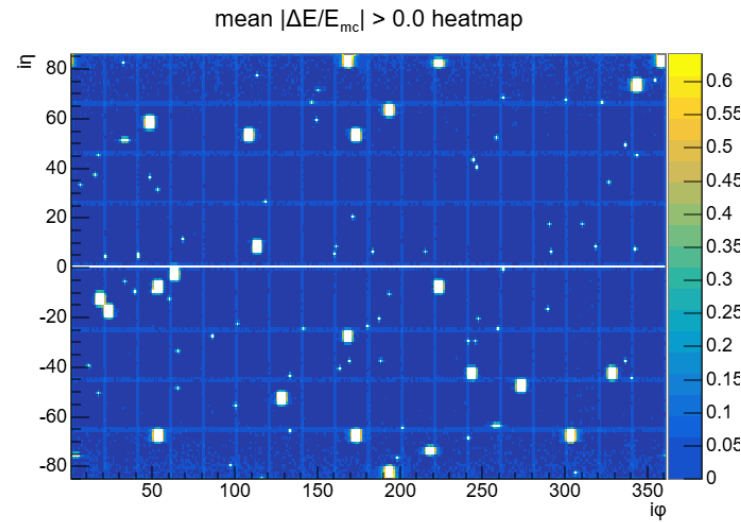
Density map of E_{reco}/E_{mc} vs E_{mc}



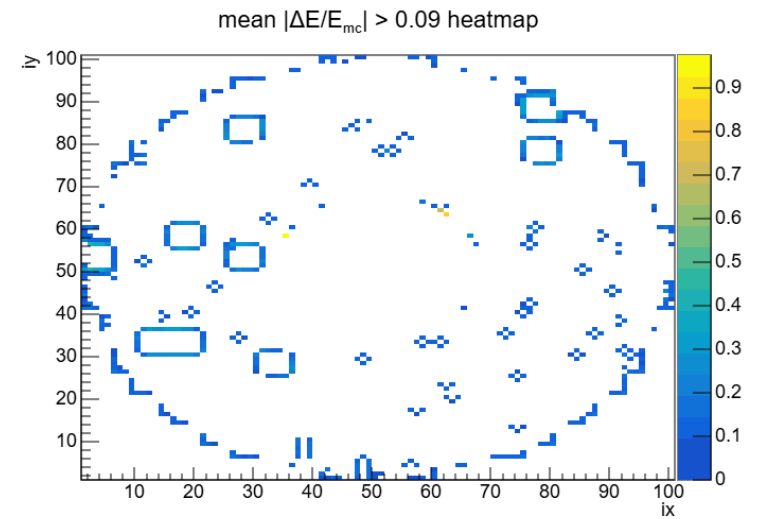
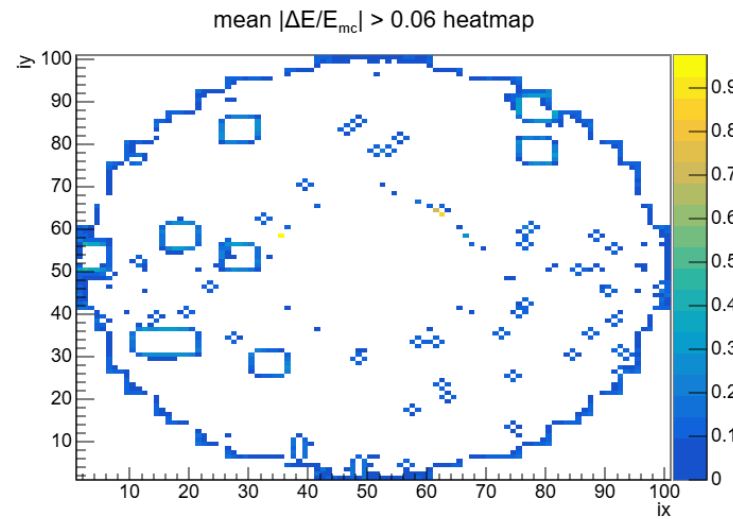
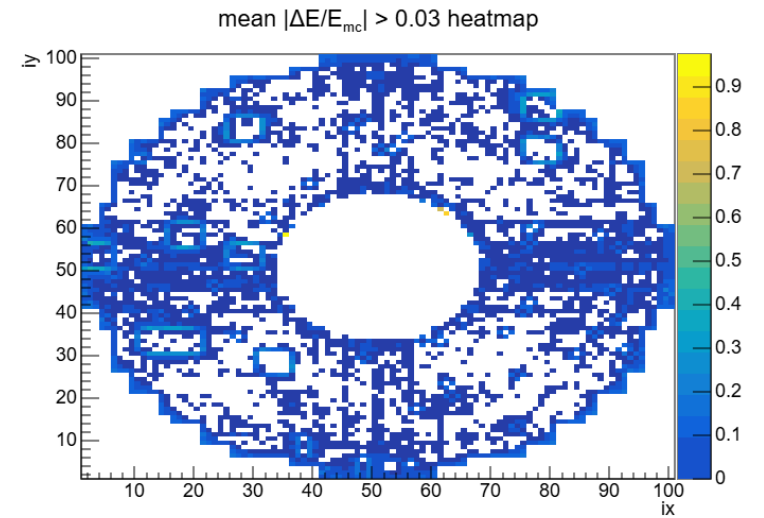
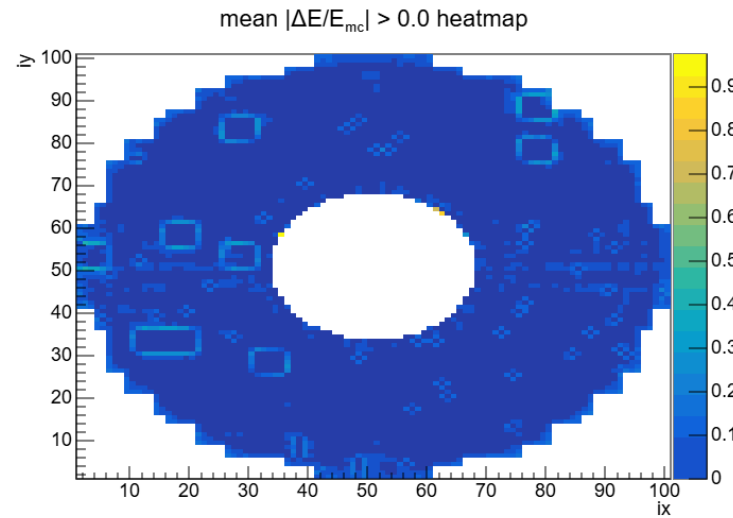
Error vs shower containment

- $(E_{5X5} - E_{MC})/E_{MC}$ plots
- 4 distinct regions
 - 1. Gaussian like distribution at higher containments
 - 2. $y = x$ line region, i.e. $E_{MC} = E_{5X5}$
 - 3. High density region near (-1,-1)
 - 4. Spread of points with high error along the x-axis, irrespective of x values
- Region 1 is what we want the ML model to focus on; thus, cuts should be applied to eliminate regions 2-4.

Barrel (EB)



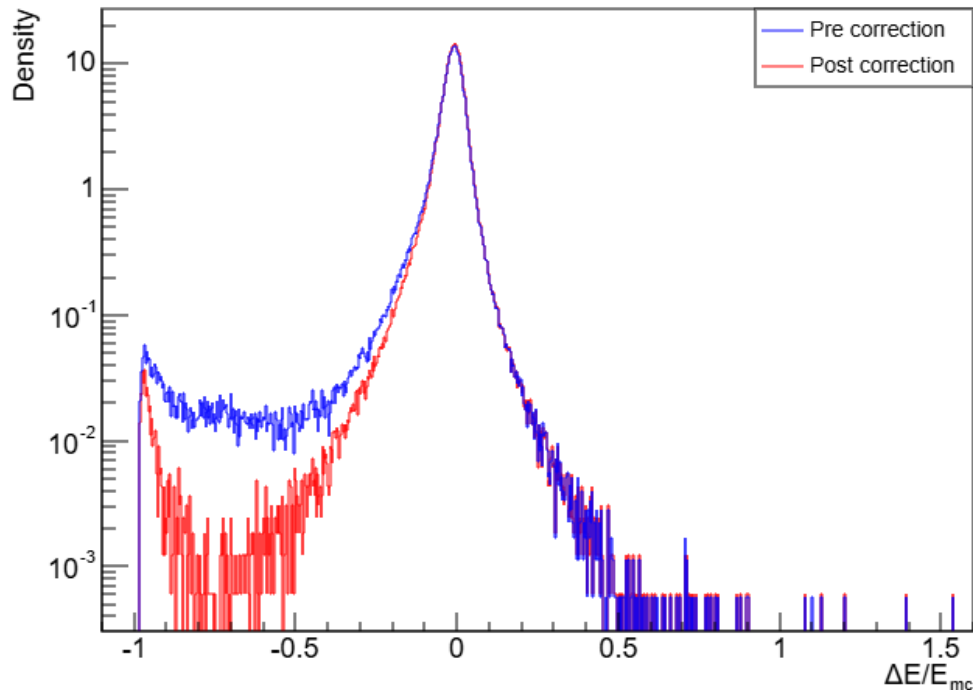
Endcaps (EE)



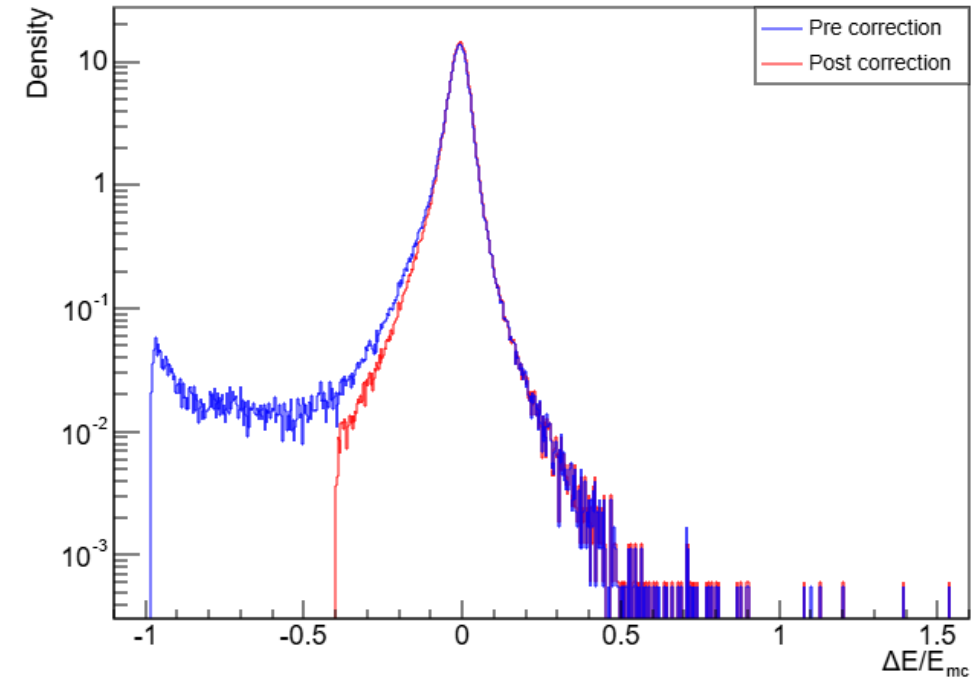
Endcap cutoffs

- Cuts in EE can be applied by:
 - Removing points near dead crystals and outer edges of the EE

$\Delta E/E_{mc}$ with and without DC cut offs | 8.0% removed



$\Delta E/E_{mc}$ with and without DC and $\Delta E/E_{mc} > -0.4$ cut offs | 8.0% removed



ML model – Data settings

- Data is separated into a training, validation and testing dataset:
 - Training: X samples
 - Validation: X samples
 - Testing: X samples
- Corrections applied:
 - Near-dead crystal pixels cutoff
 - Super-module transition in EB and outer EE edge cutoff
 - Sharp cutoff of: $-0.4 < \Delta E / E_{MC}$
- Separate models created for EE and EB

XGBoost (BDT)

- List of input features:

- Tracker:

- detain
 - dphiin
 - fbrem
 - ooemoop
 - trkpt
 - trketa
 - trkphi
 - trkq

- ECAL:

- e1x5
 - e2nd
 - e2x5B, e2x5L, e2x5M,
 - e2x5R, e2x5T
 - e5x5
 - eB, eL, eR, eT
 - eMax
 - eta ieta
 - iphi phi
 - pse
 - pt
 - r9
 - sieie

- Labels:

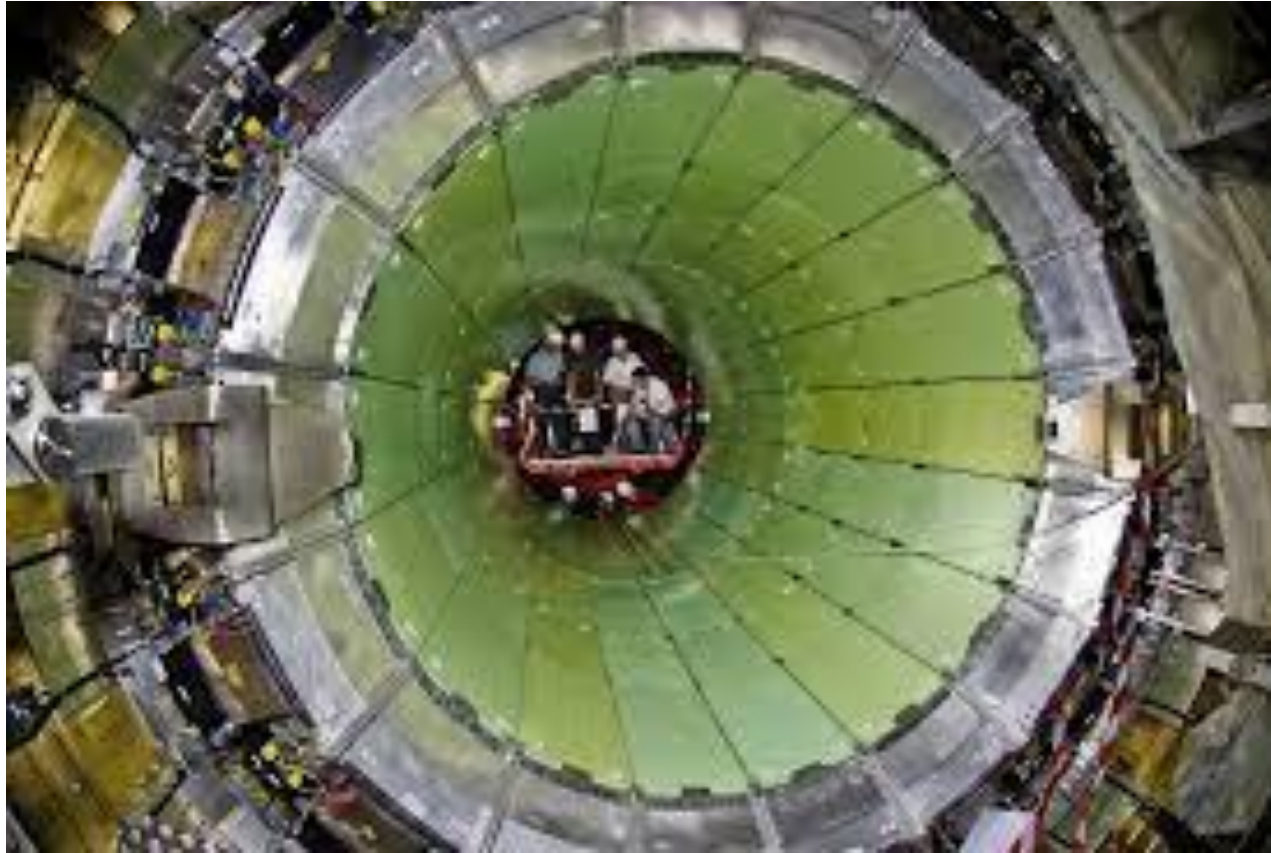
- $y_{true} = E_{MC}/E_{reco}$

- Loss:

- $= (E_{pred} - E_{MC})/E_{MC})^2$

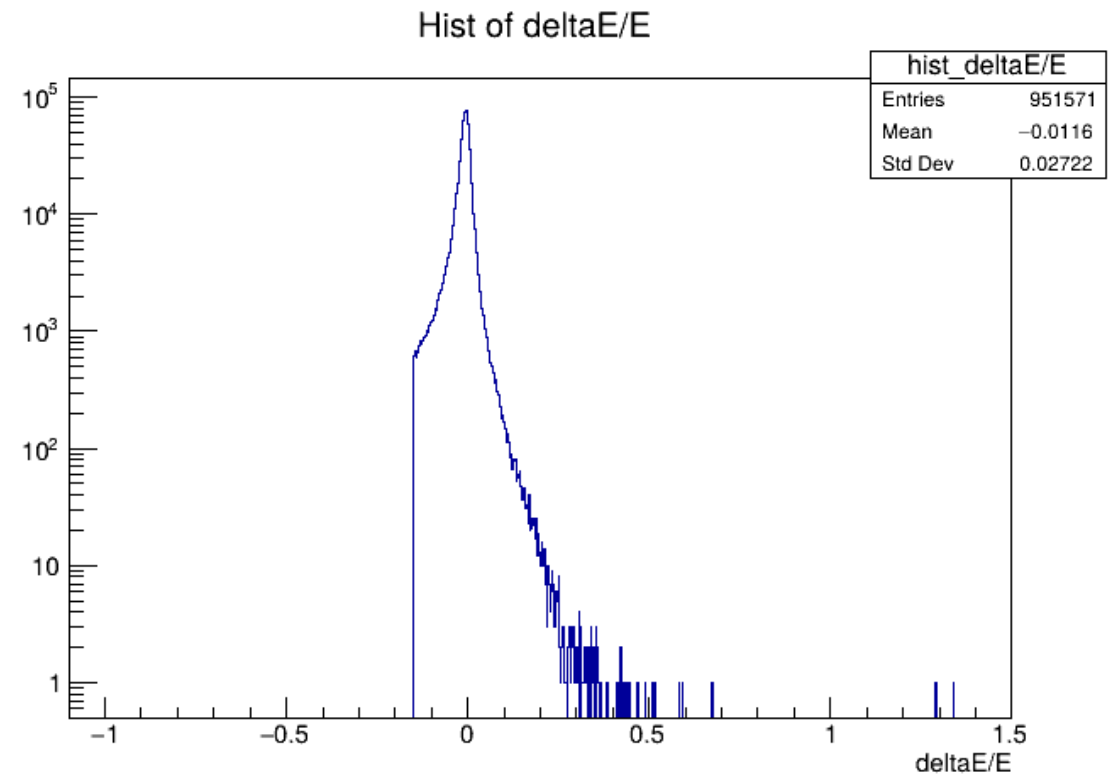
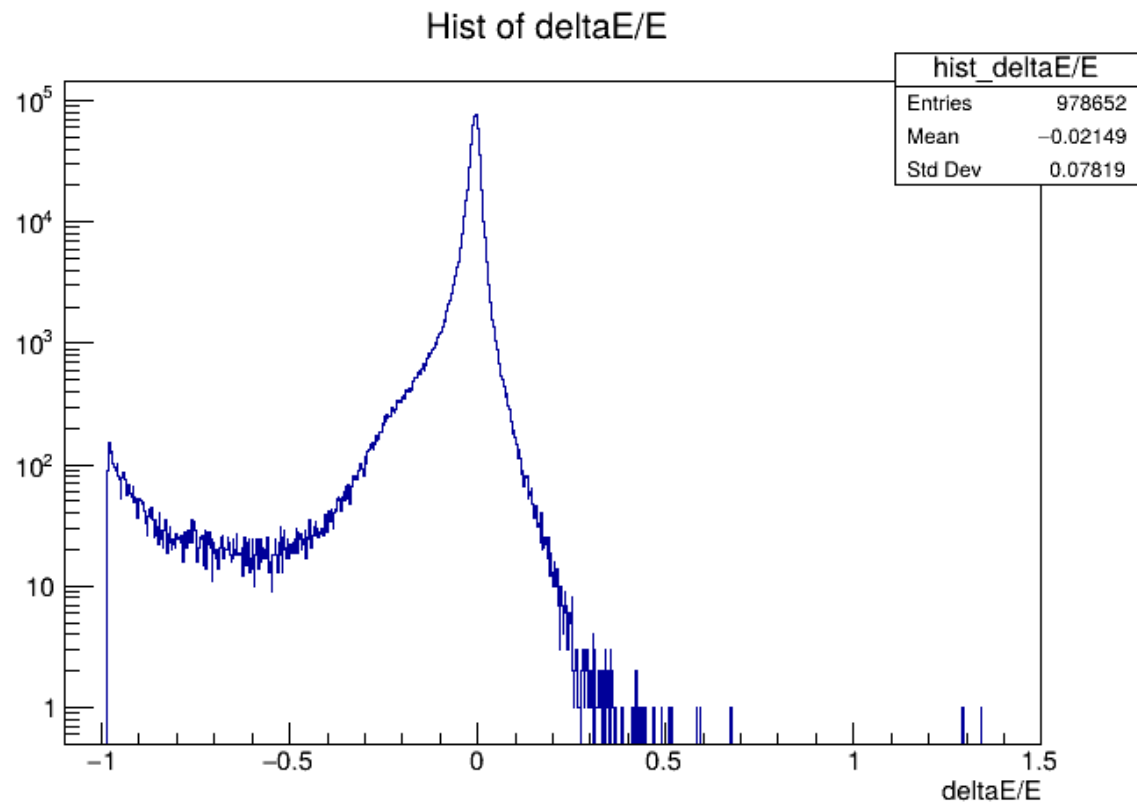
- $= (\frac{y_{pred}}{y_{true}} - 1)^2$

- No normalization (not necessary for BDT)

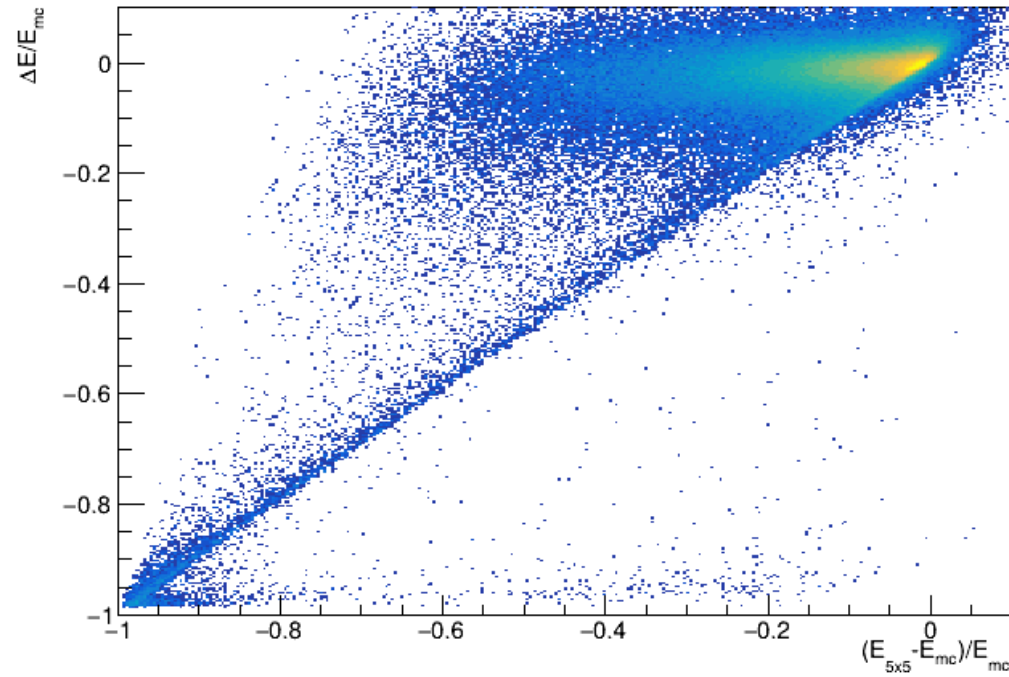
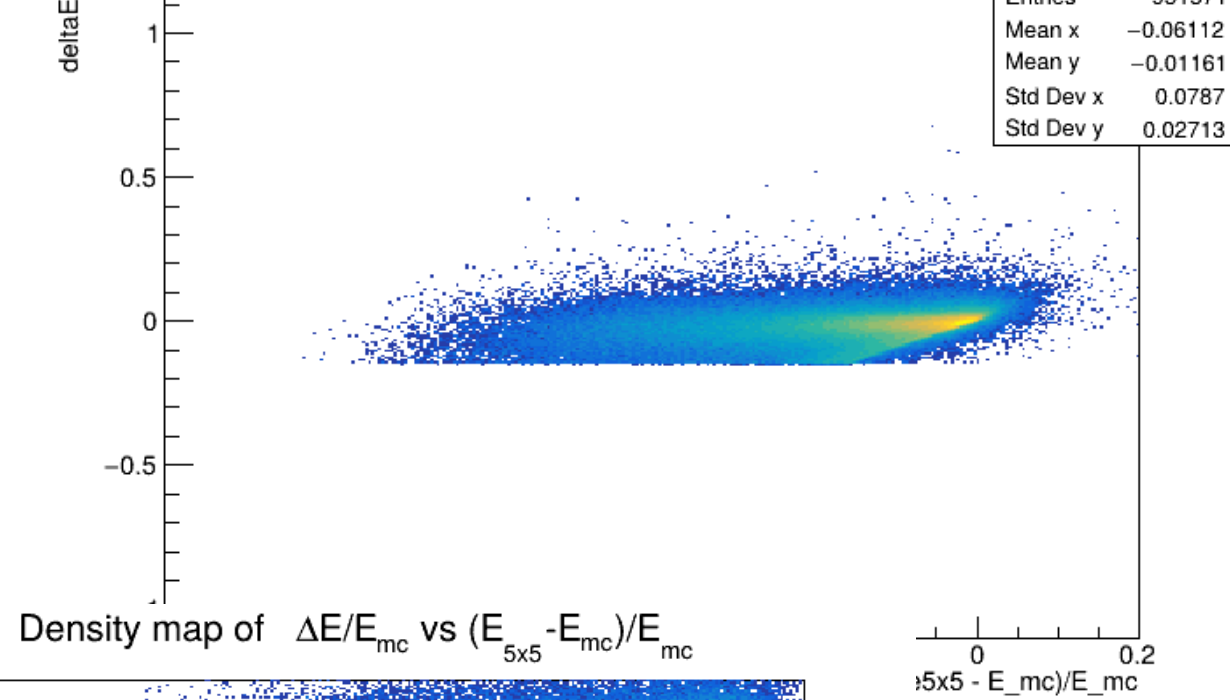
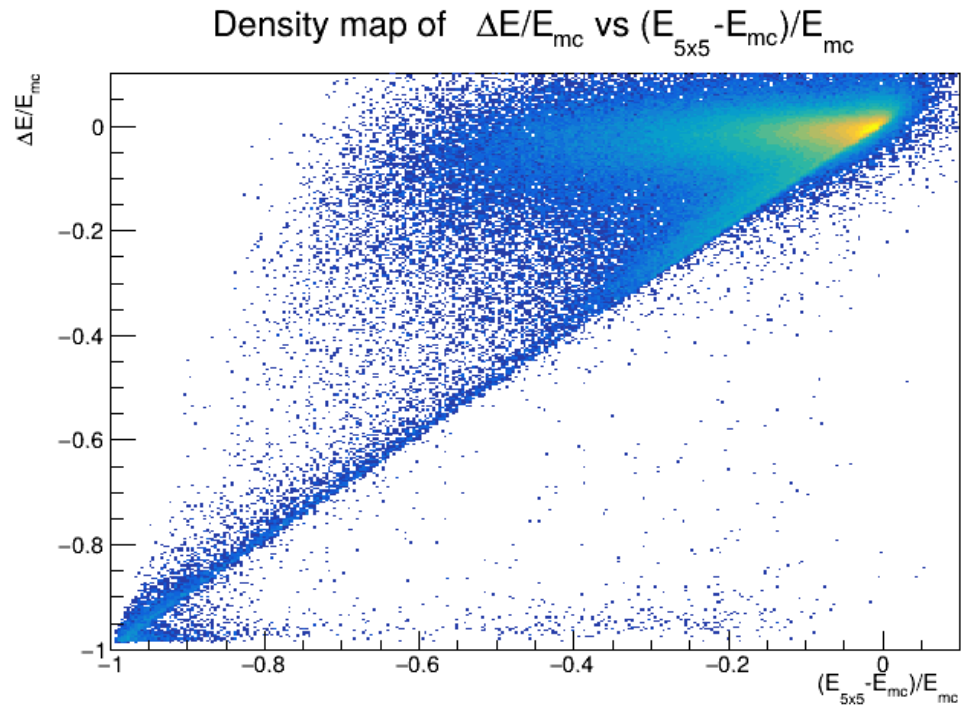


BACKUPS ON BACKUPS

Pre vs post corrections

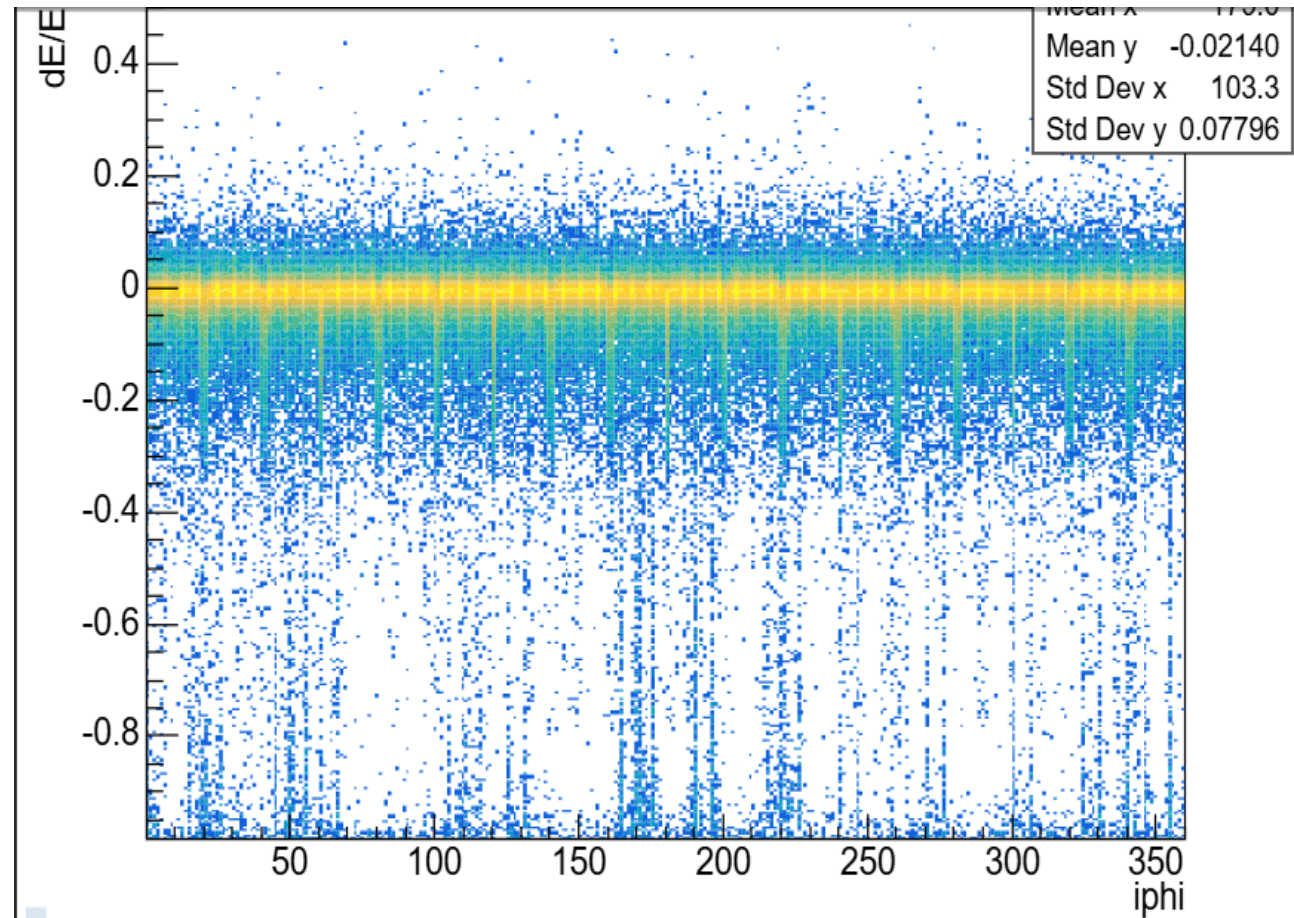


Pre vs post corrections

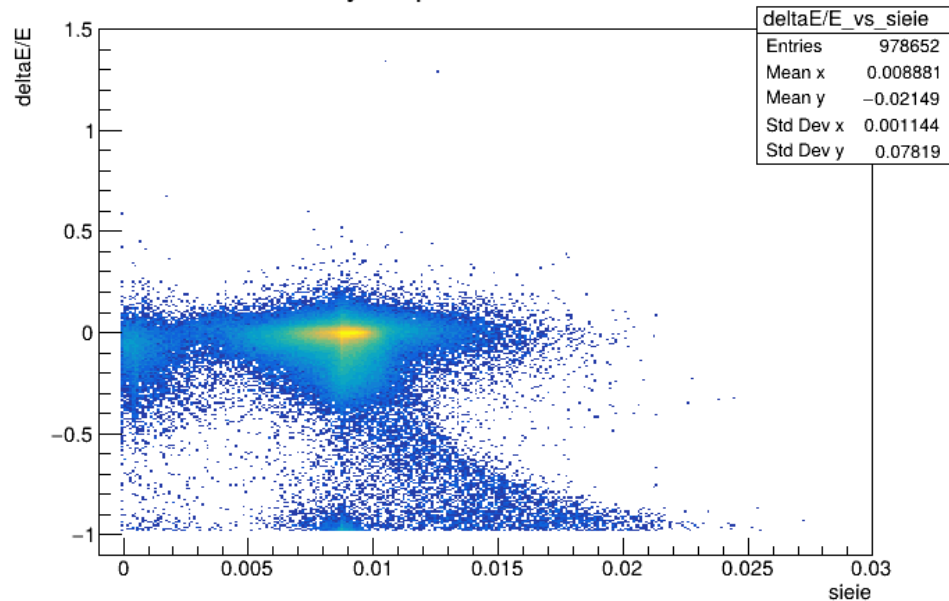


Investigating errors

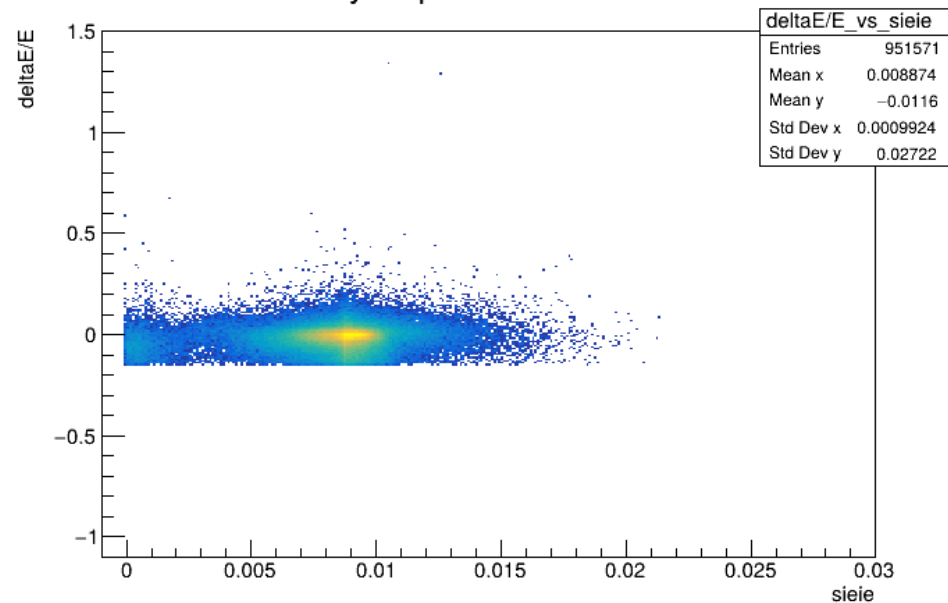
- Plot of err vs iphi:
 - 18 Fringes – 18 Super modules



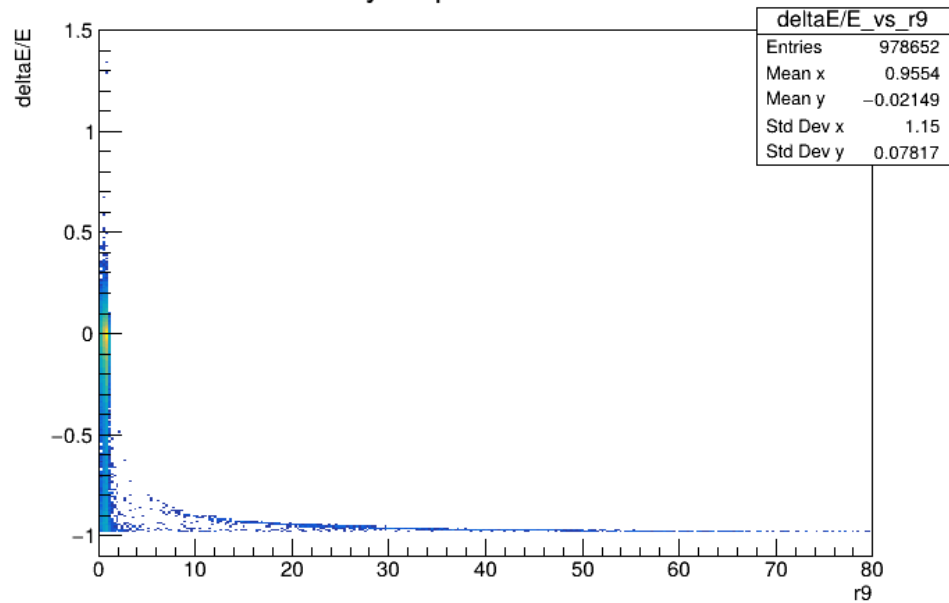
Density map of deltaE/E vs sieie



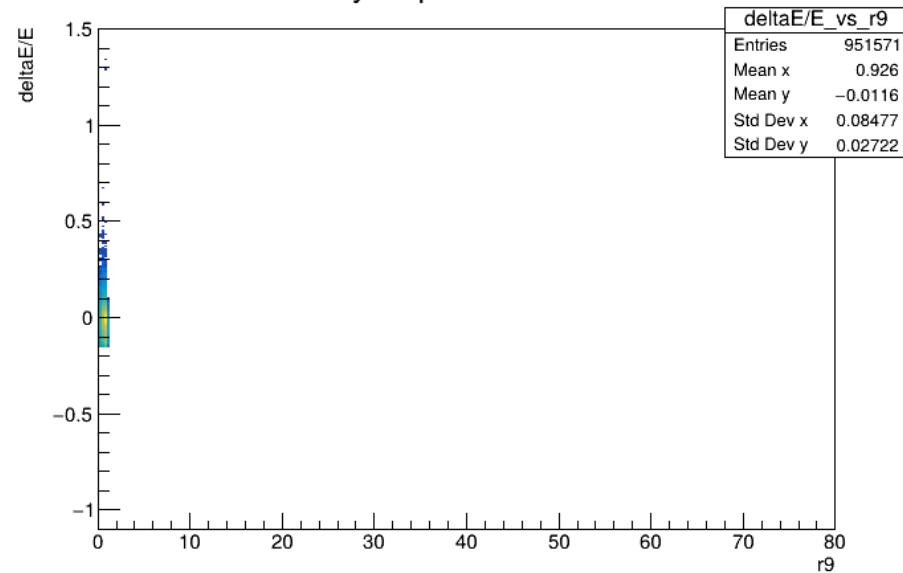
Density map of deltaE/E vs sieie



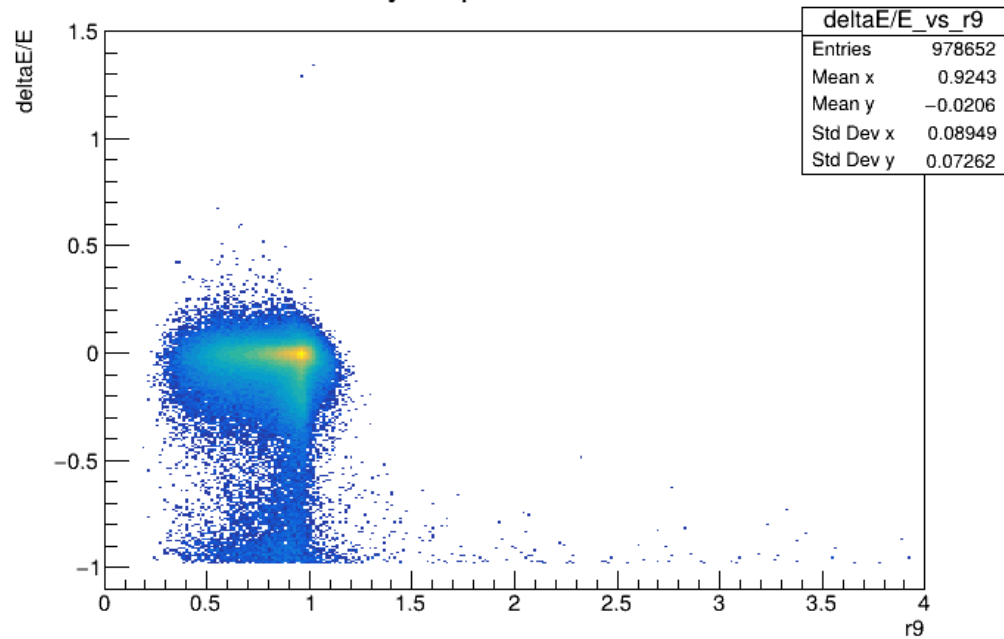
Density map of deltaE/E vs r9



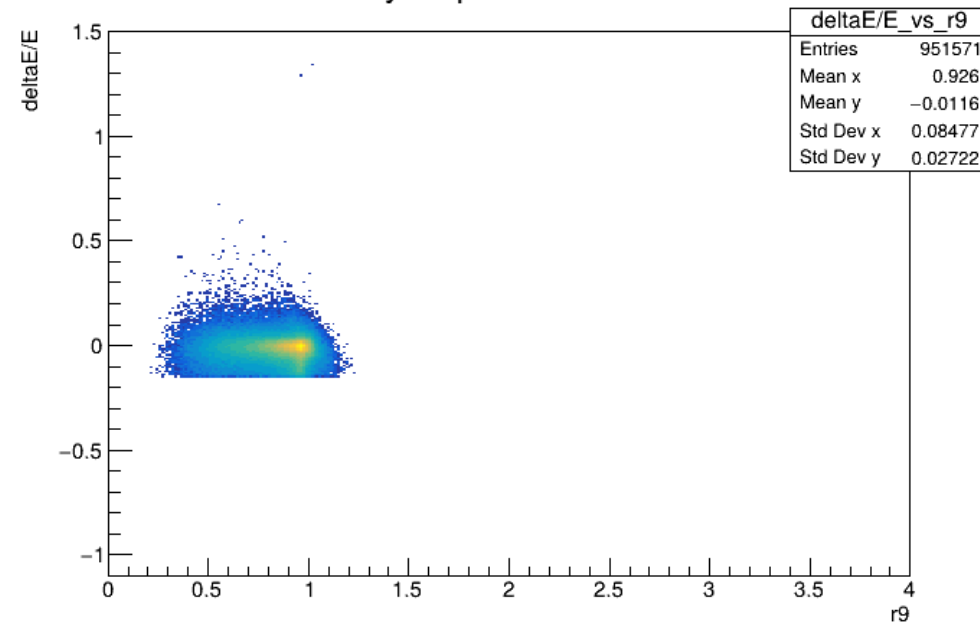
Density map of deltaE/E vs r9



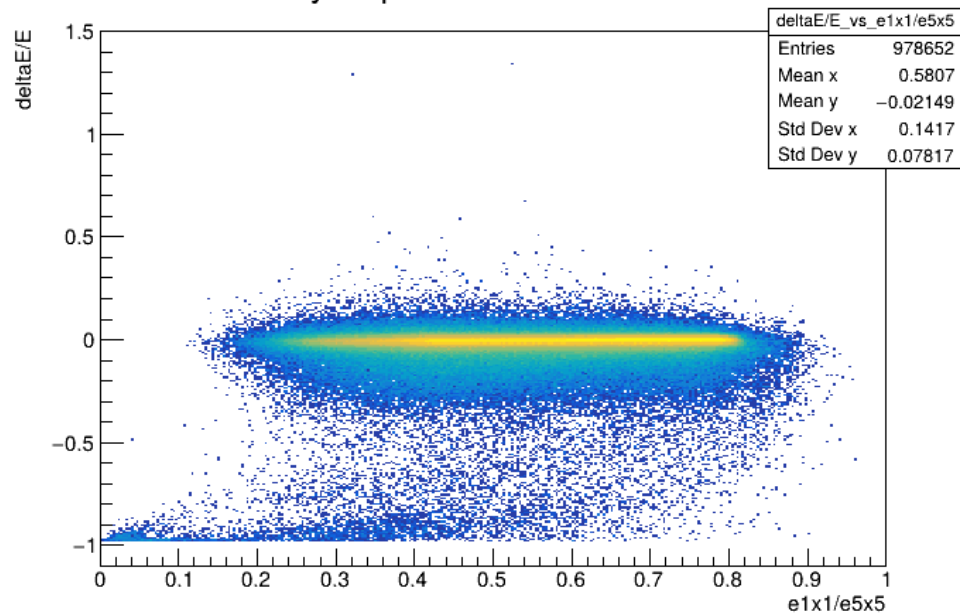
Density map of deltaE/E vs r9



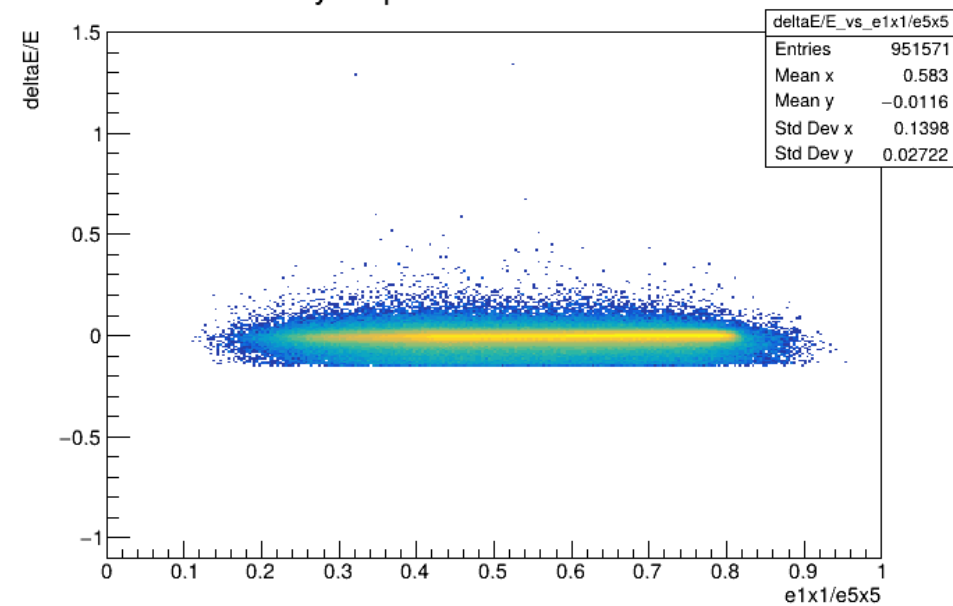
Density map of deltaE/E vs r9

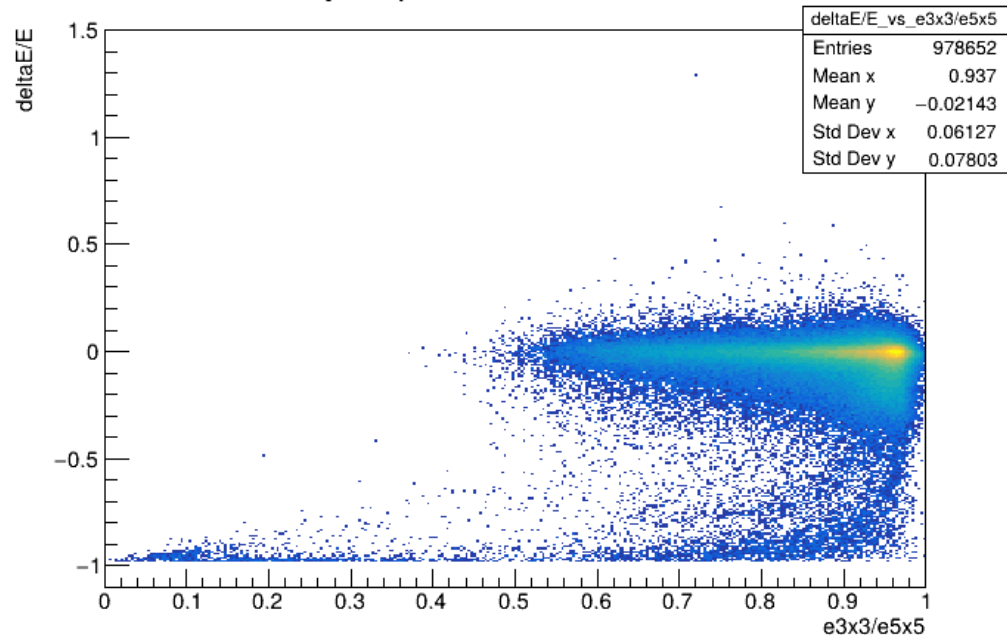
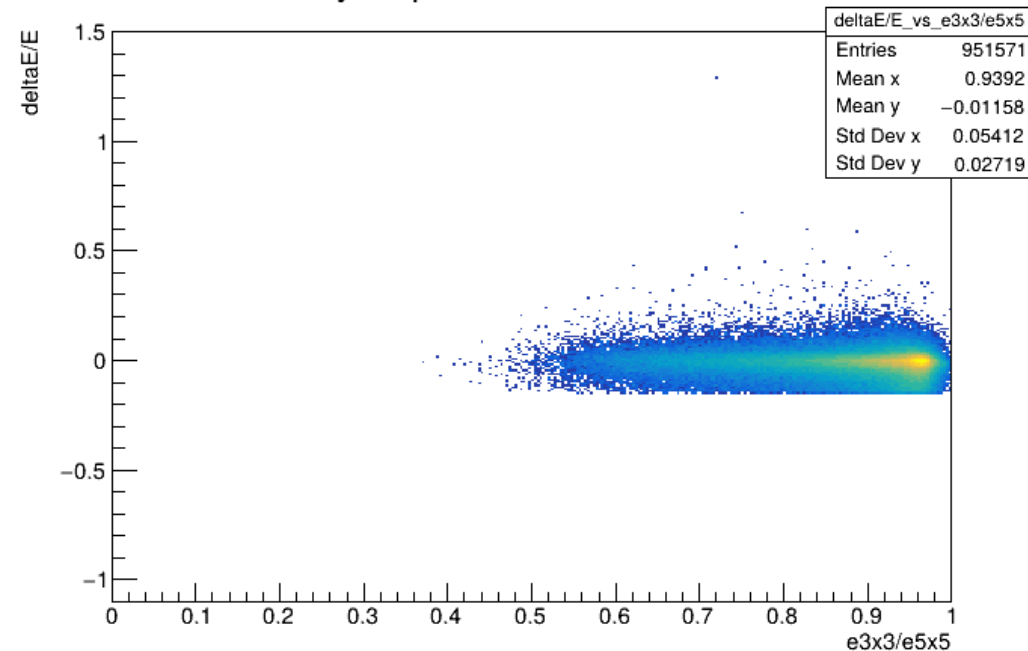
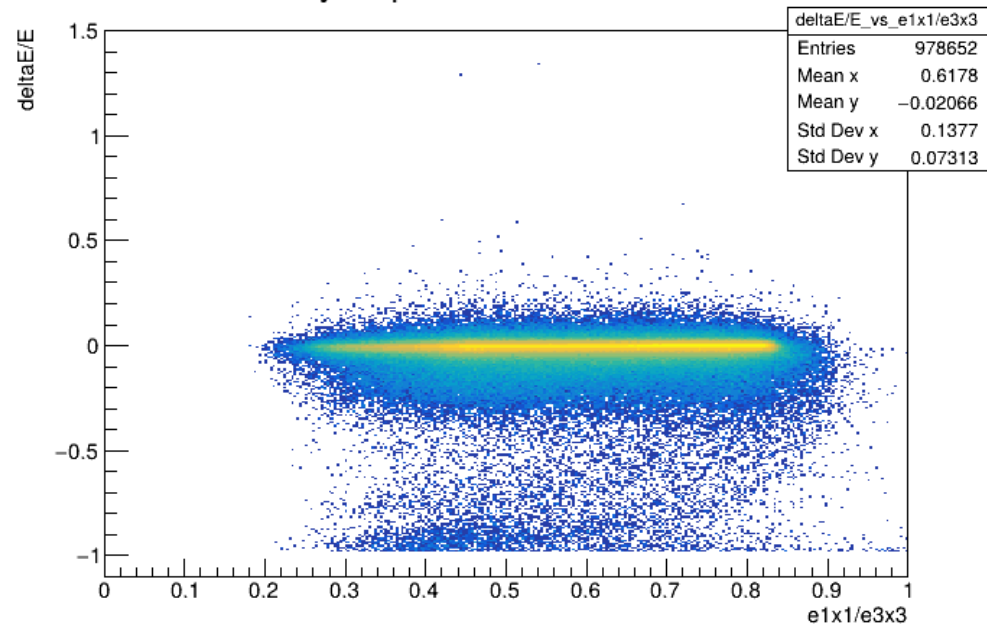
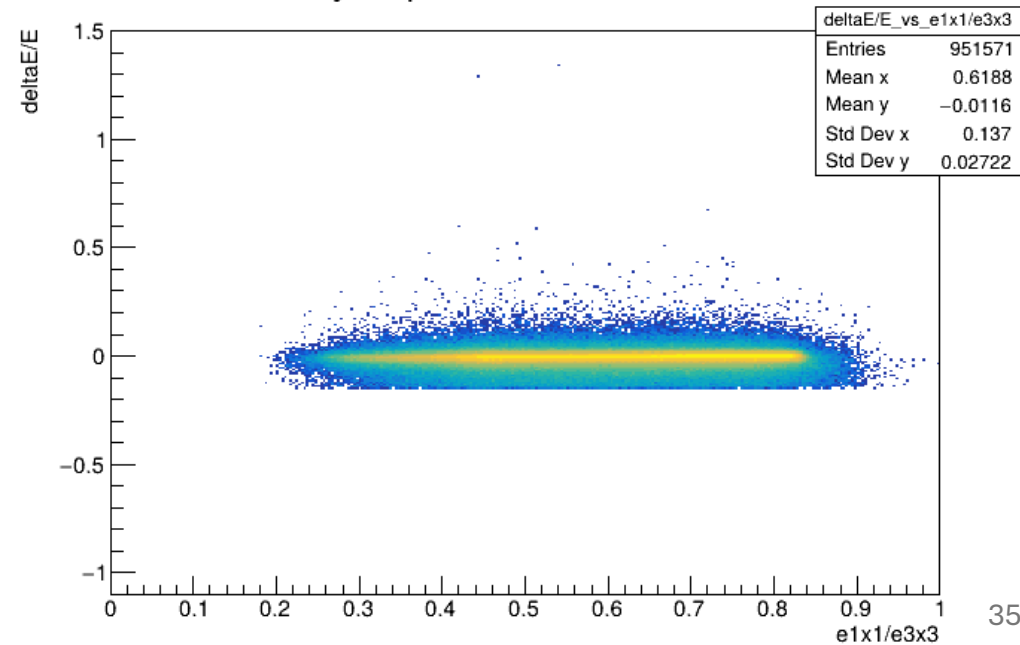


Density map of deltaE/E vs e1x1/e5x5



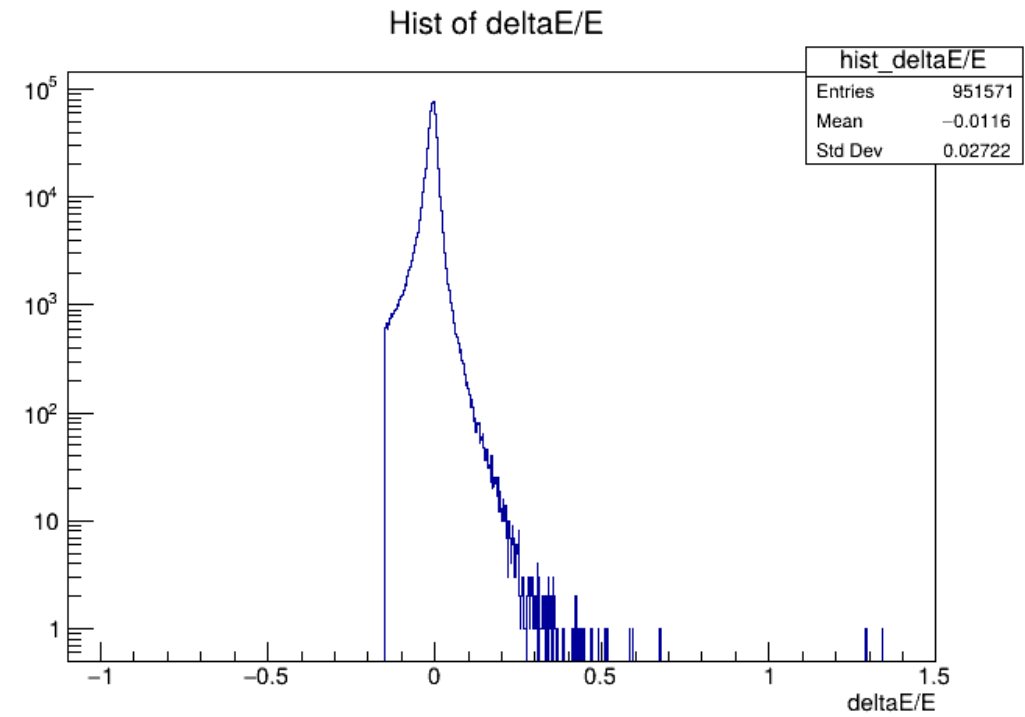
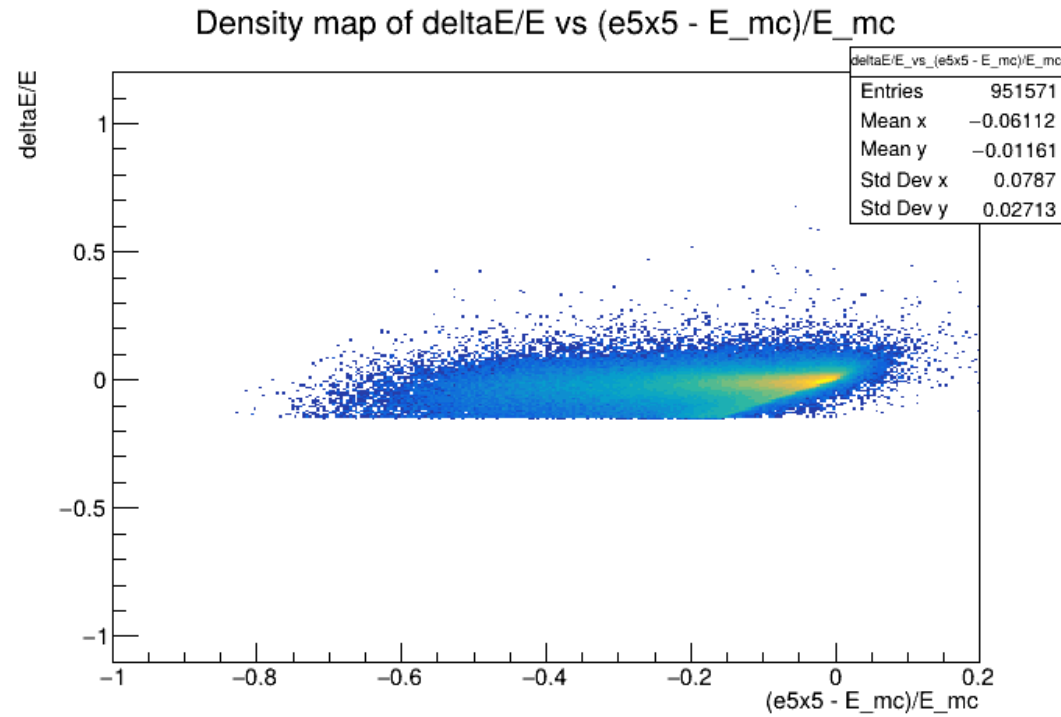
Density map of deltaE/E vs e1x1/e5x5



Density map of $\Delta E/E$ vs $e3x3/e5x5$ Density map of $\Delta E/E$ vs $e3x3/e5x5$ Density map of $\Delta E/E$ vs $e1x1/e3x3$ Density map of $\Delta E/E$ vs $e1x1/e3x3$ 

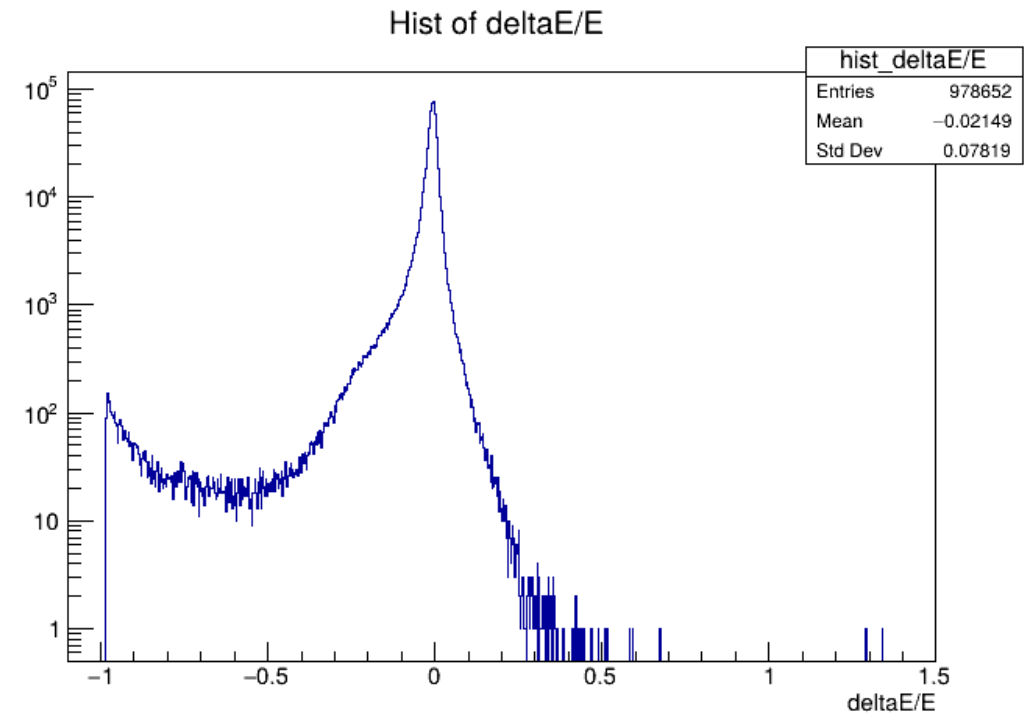
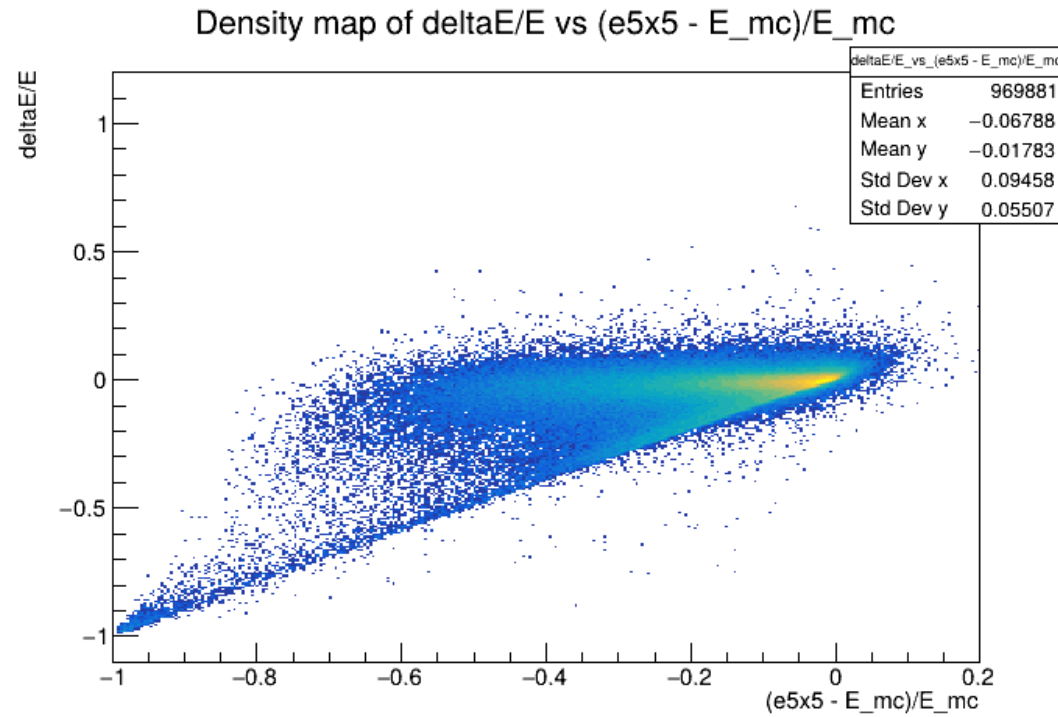
Post Correction 1:

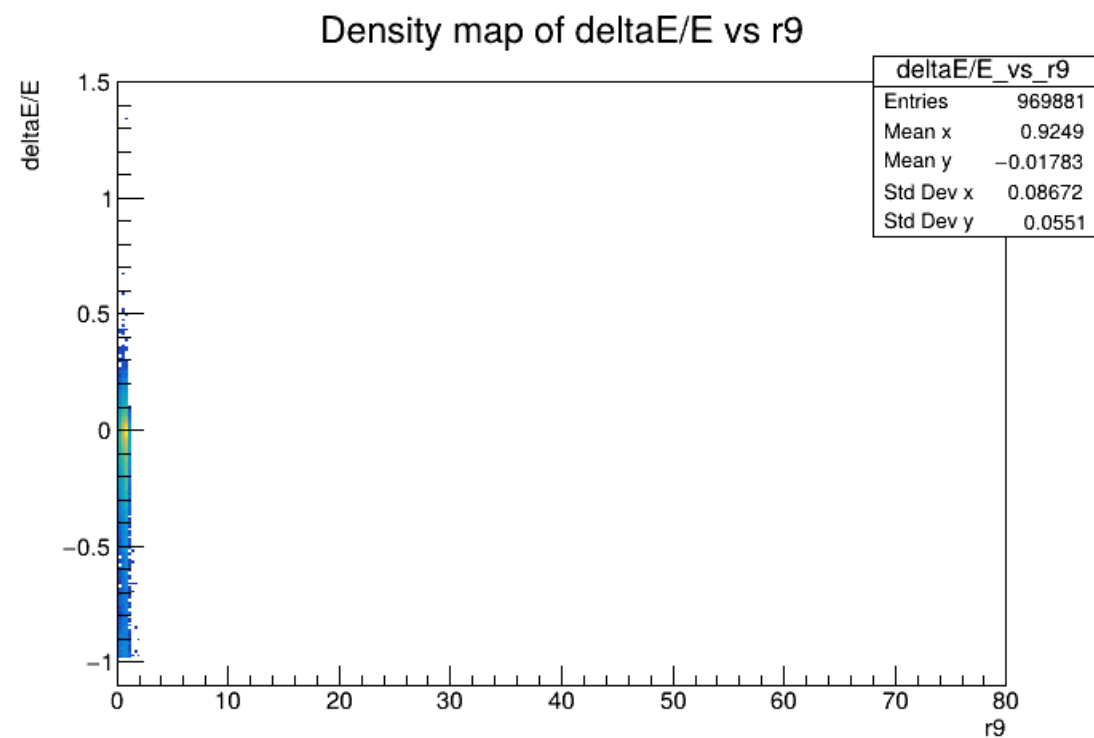
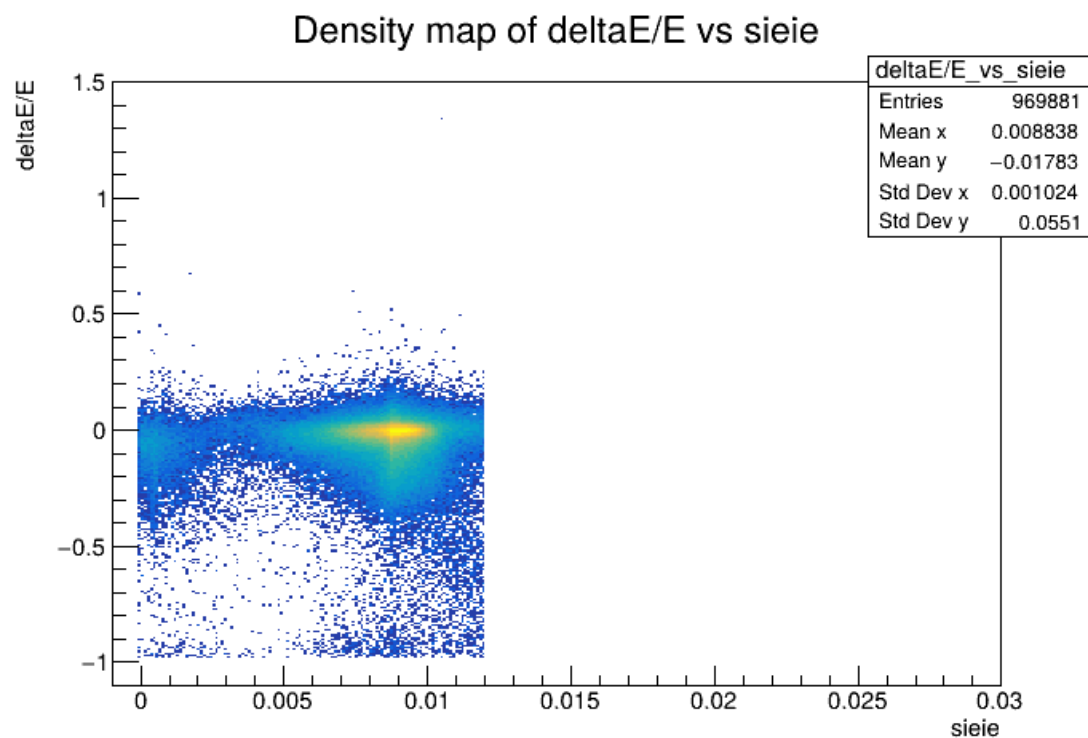
- $E_{\text{reco}} > 0.85 * E_{\text{mc}}$

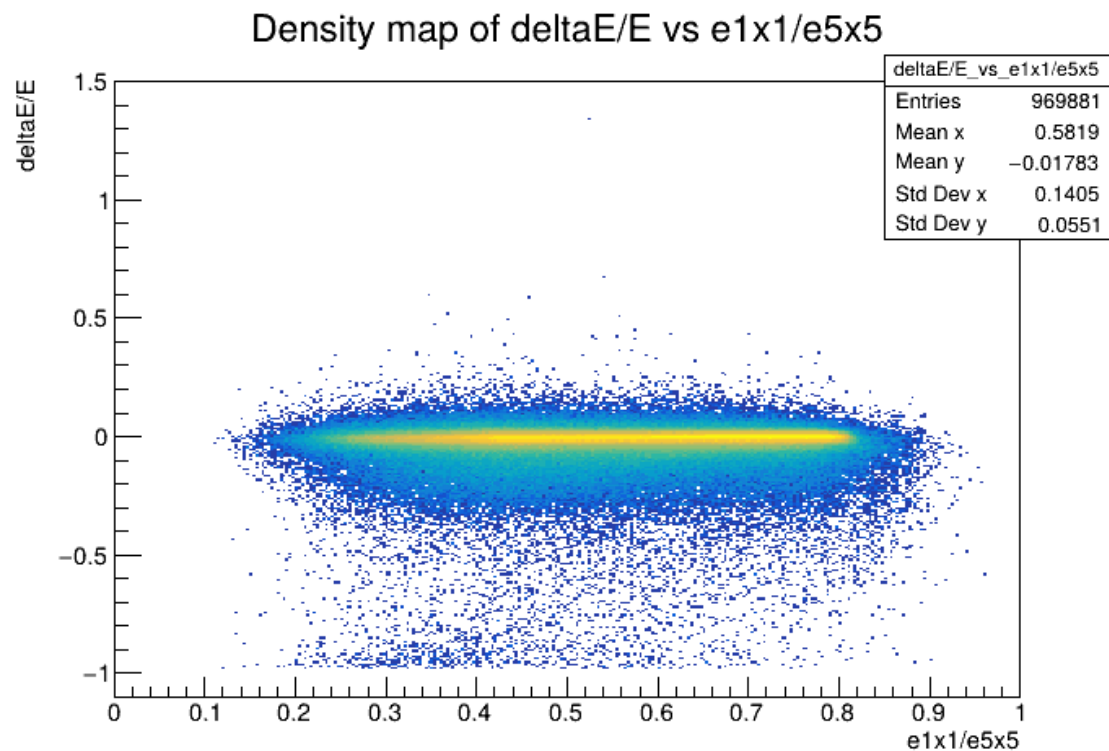
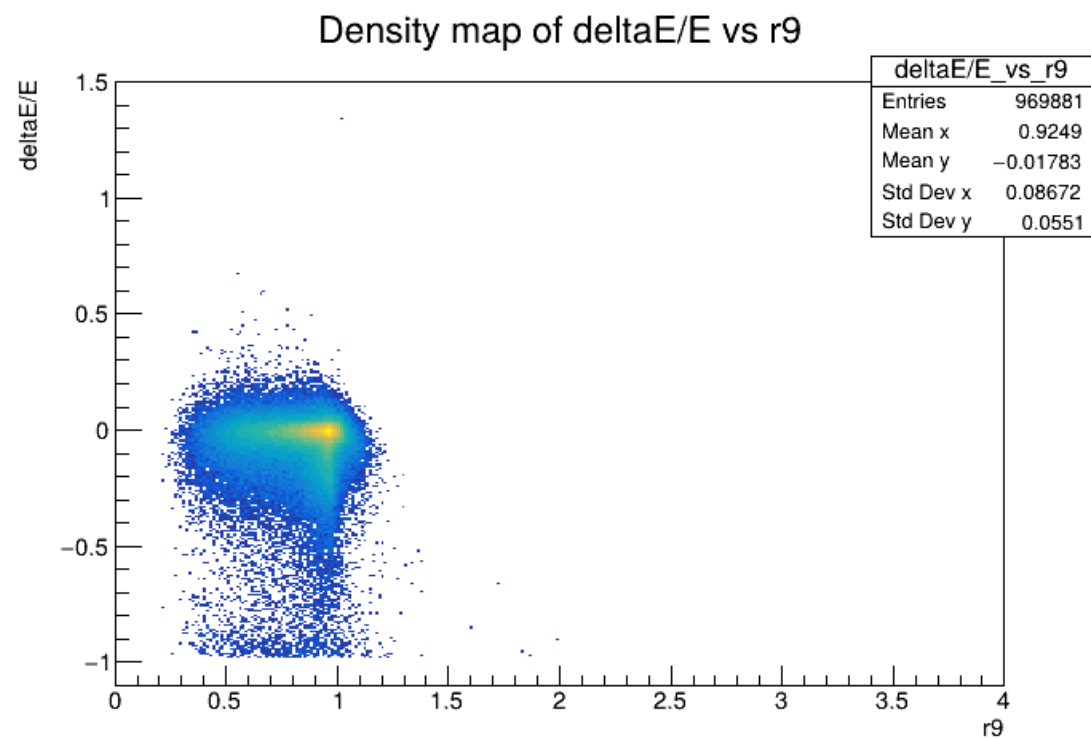


Post Correction 2:

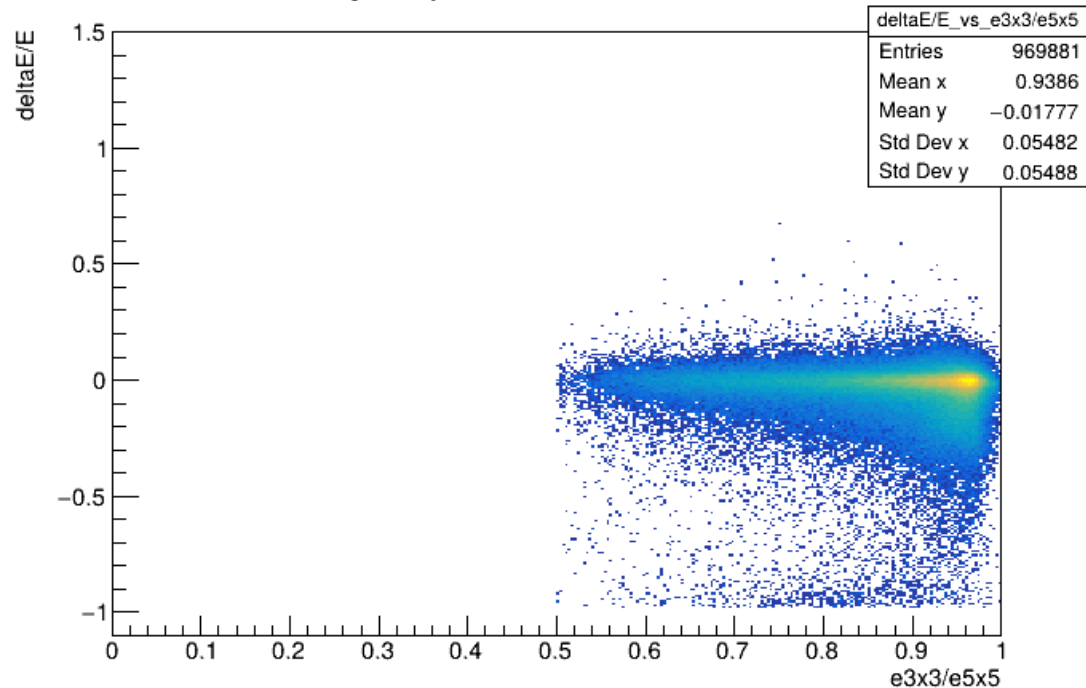
- $e3x3/e5x5 > 0.5$,
- $e_{max}/e3x3 < 1$,
- $r9 < 2$,
- $sieie < 0.012$



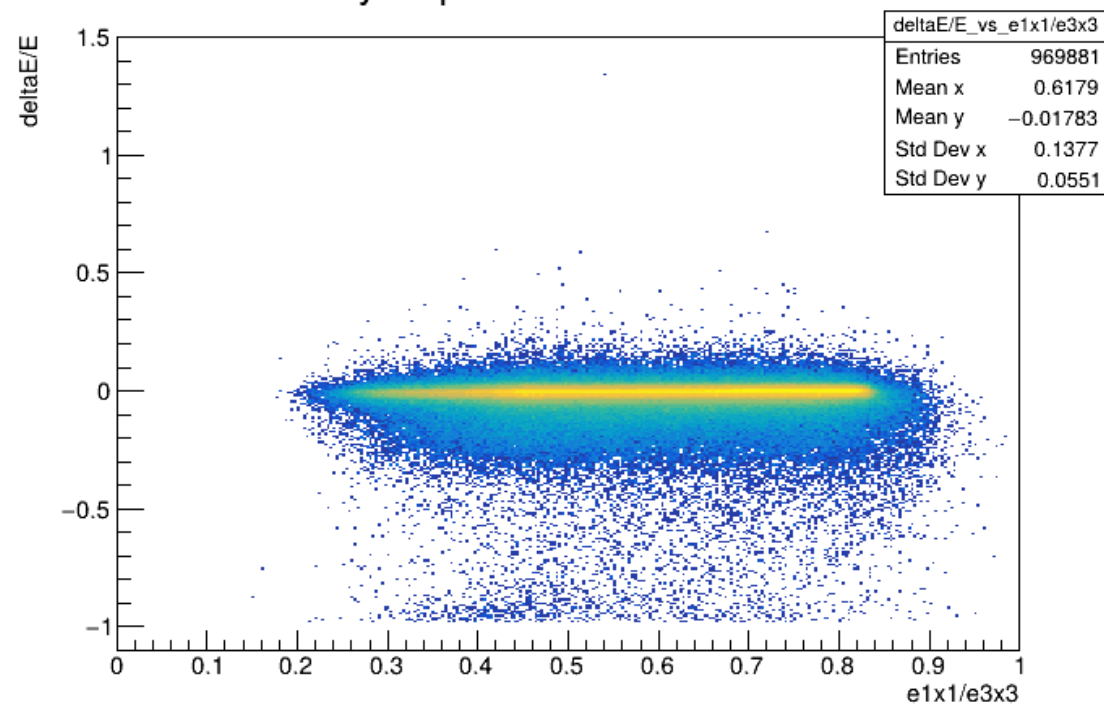


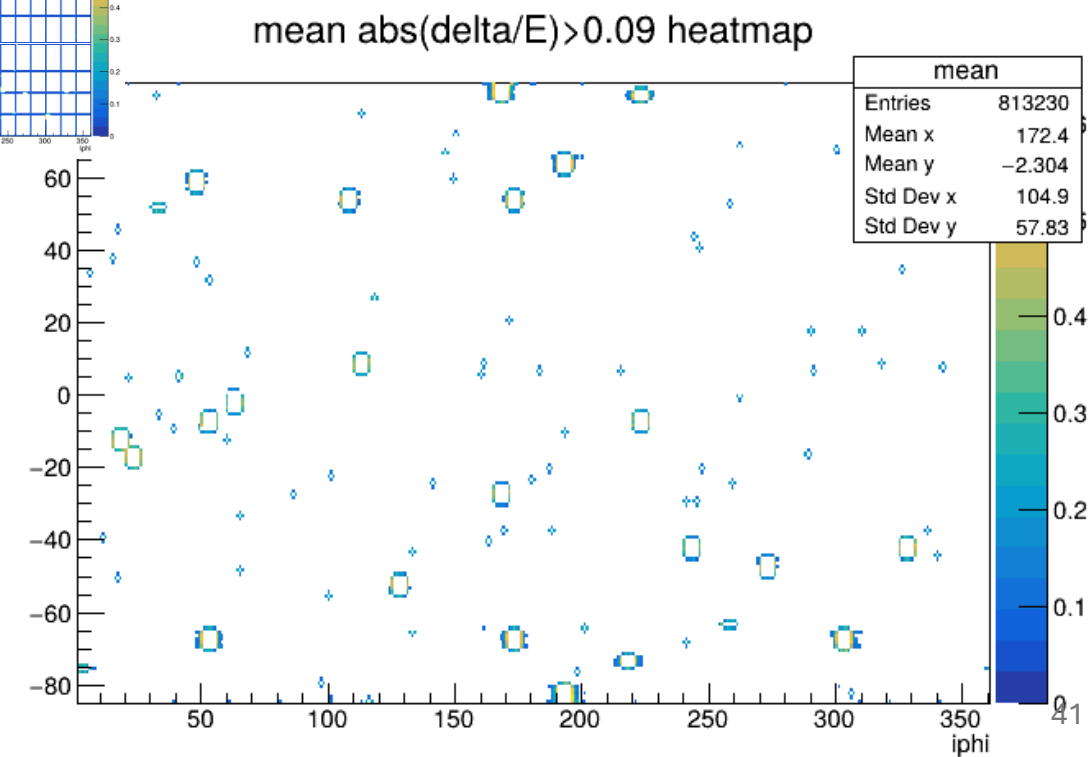
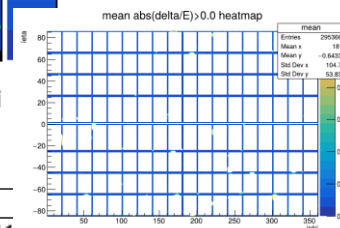
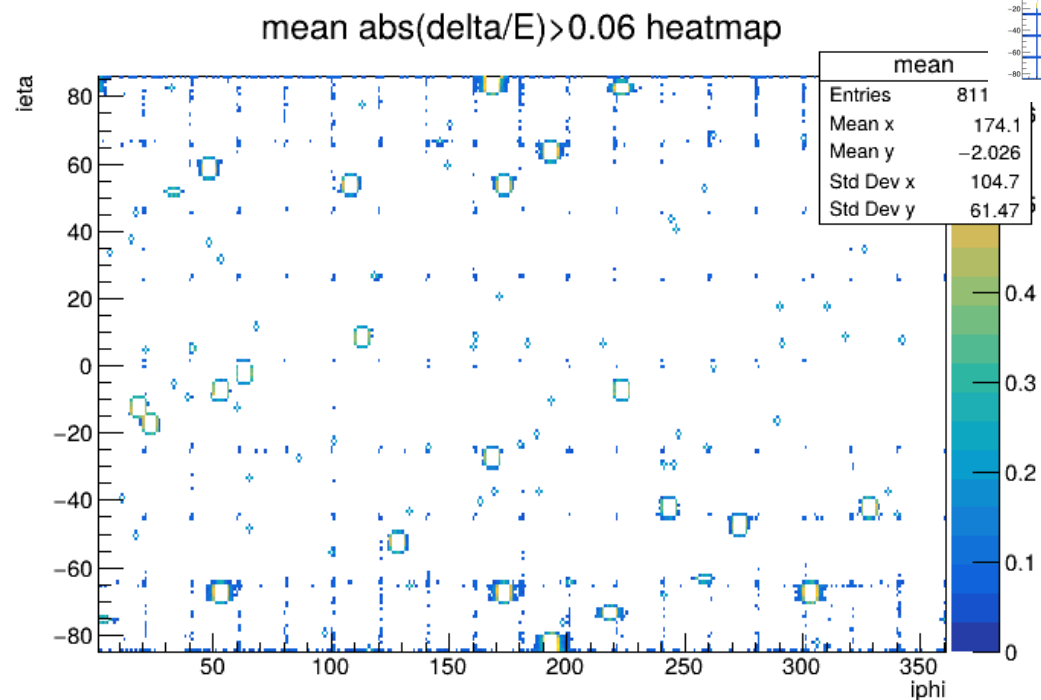
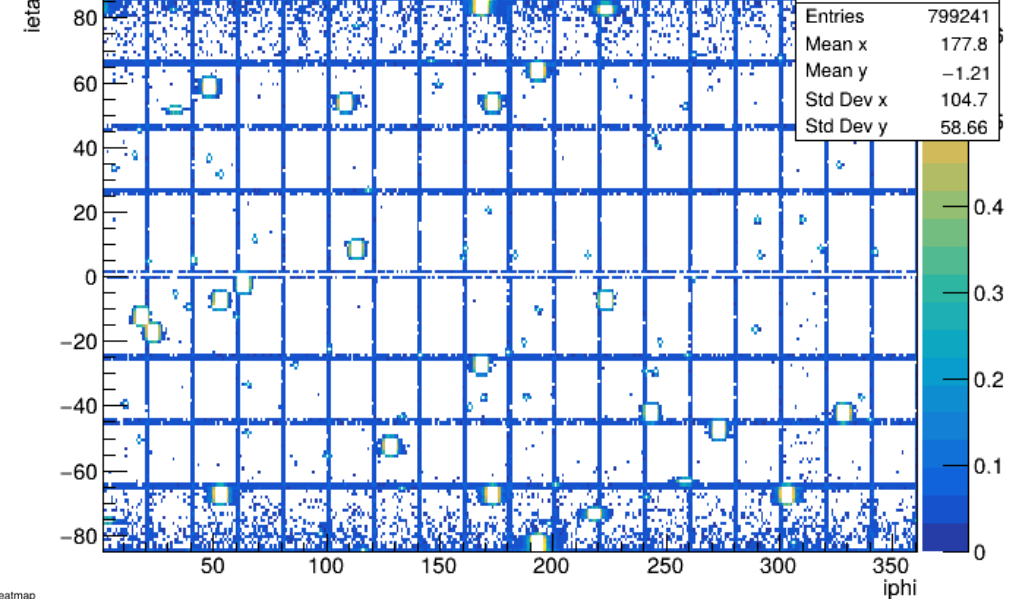
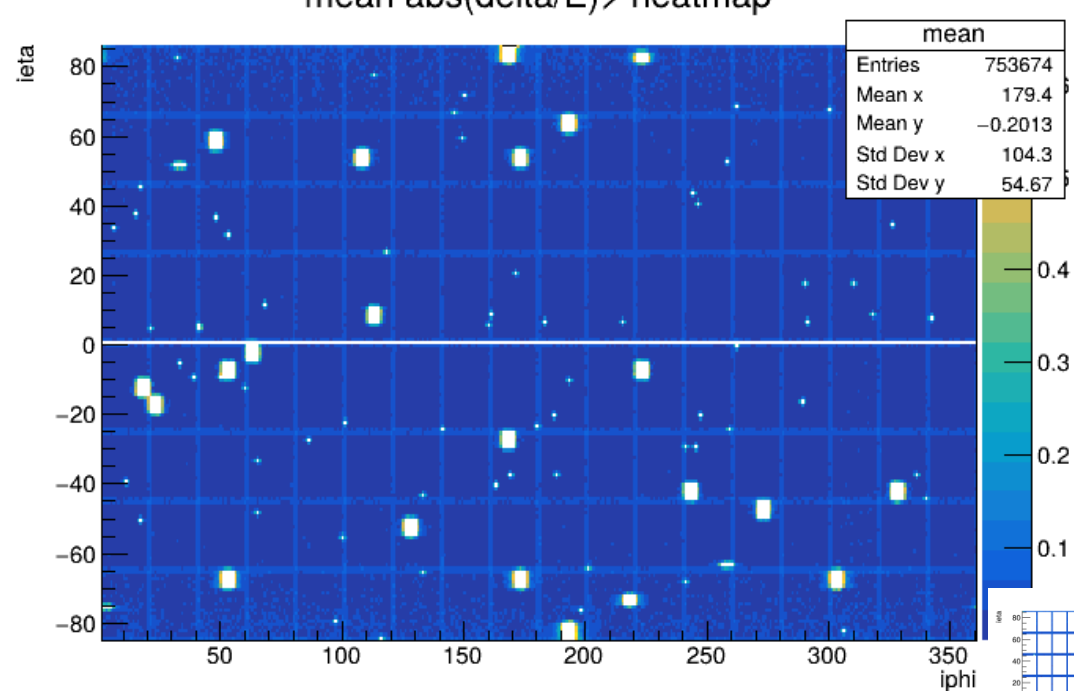


Density map of $\Delta E/E$ vs e_{3x3}/e_{5x5}

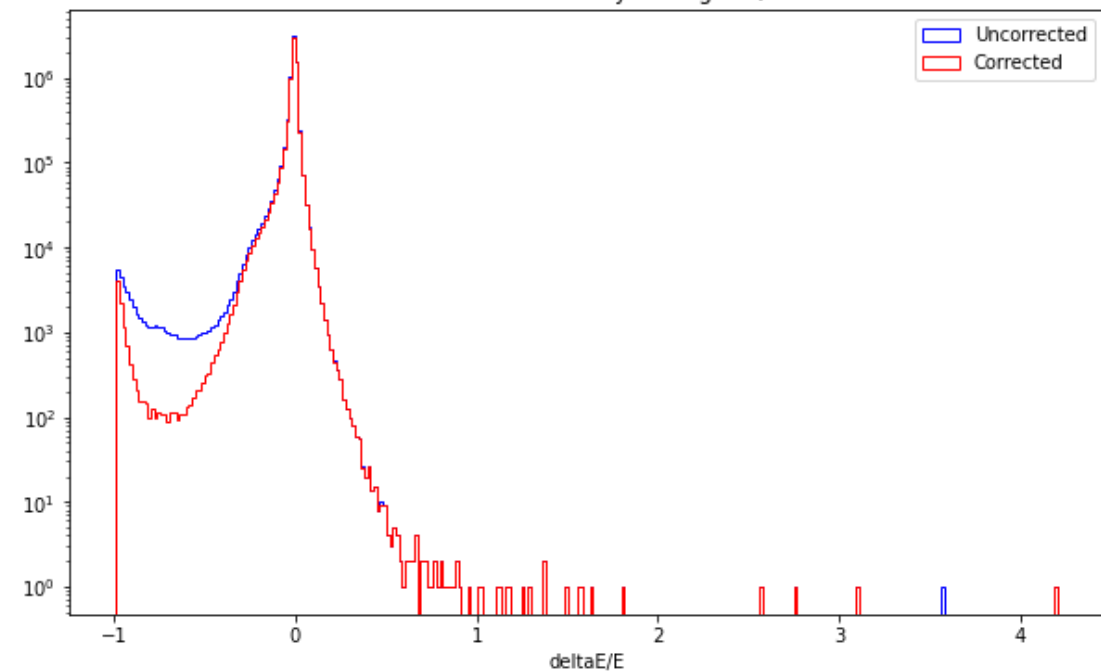


Density map of $\Delta E/E$ vs e_{1x1}/e_{3x3}

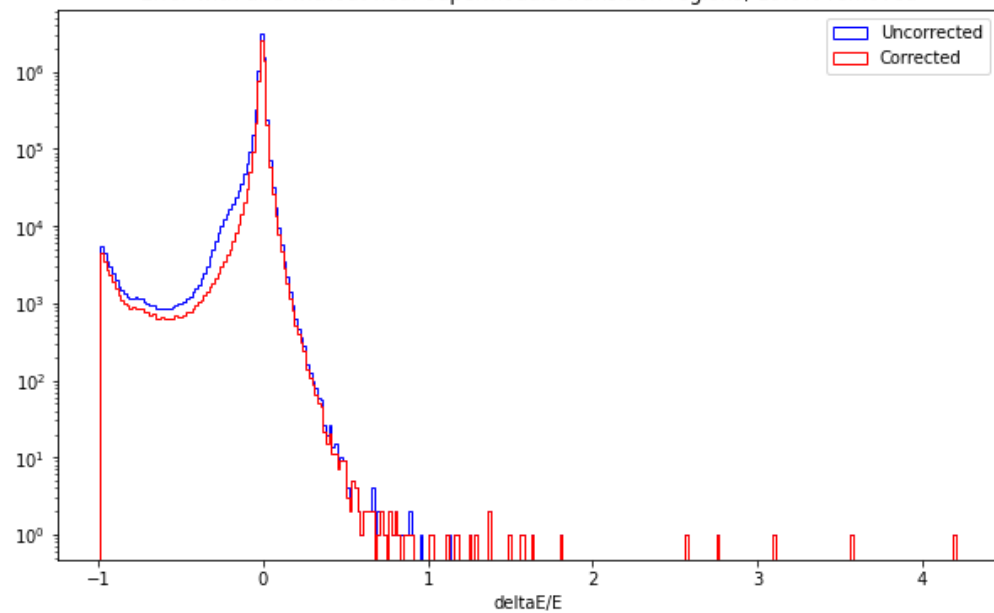




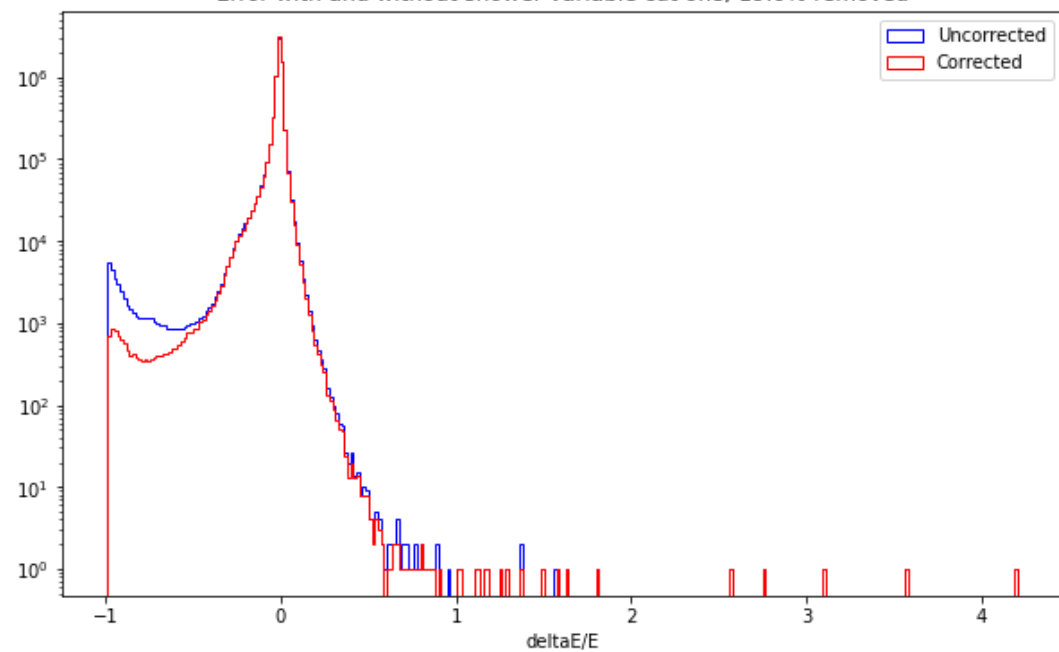
Error with and without near-deadcrystal regions; 1.94% removed



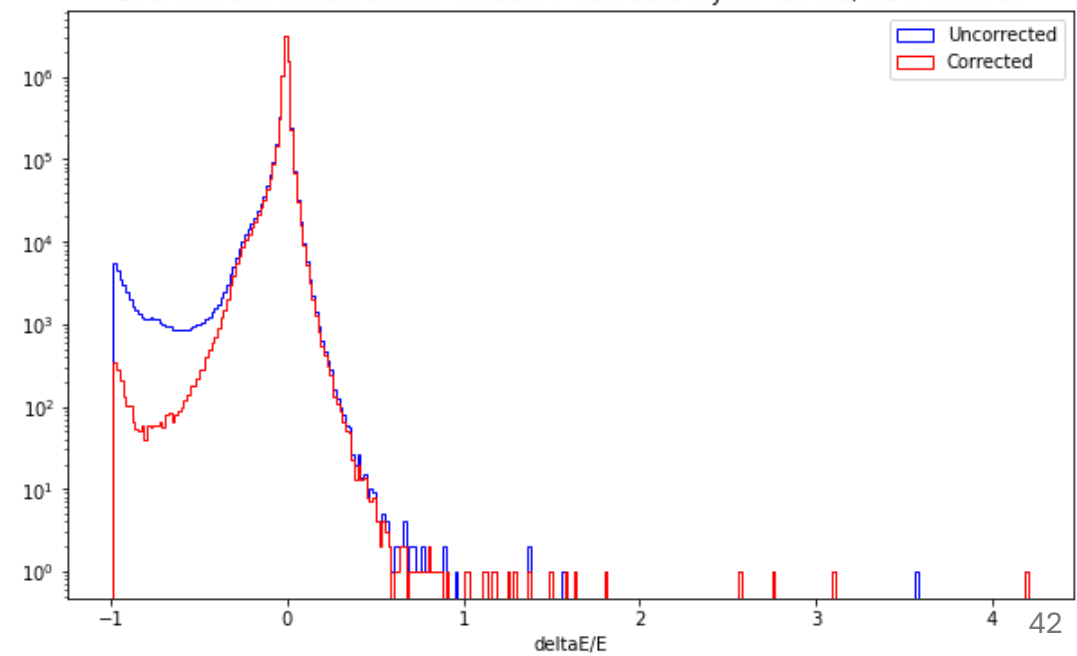
Error with and without near-supermodule transition regions; 19.0% removed



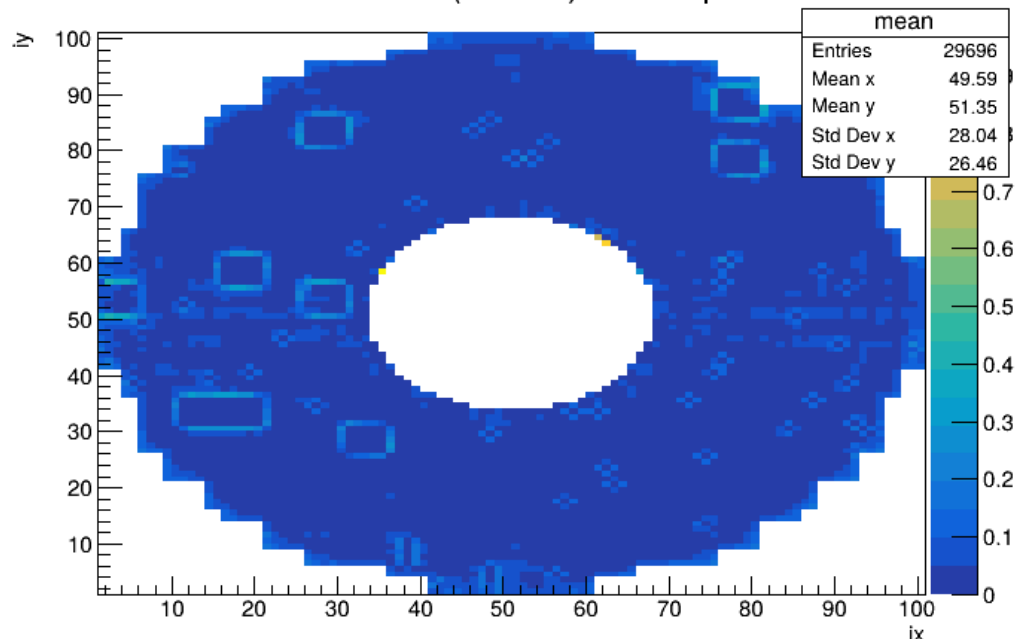
Error with and without shower variable cut offs; 19.0% removed



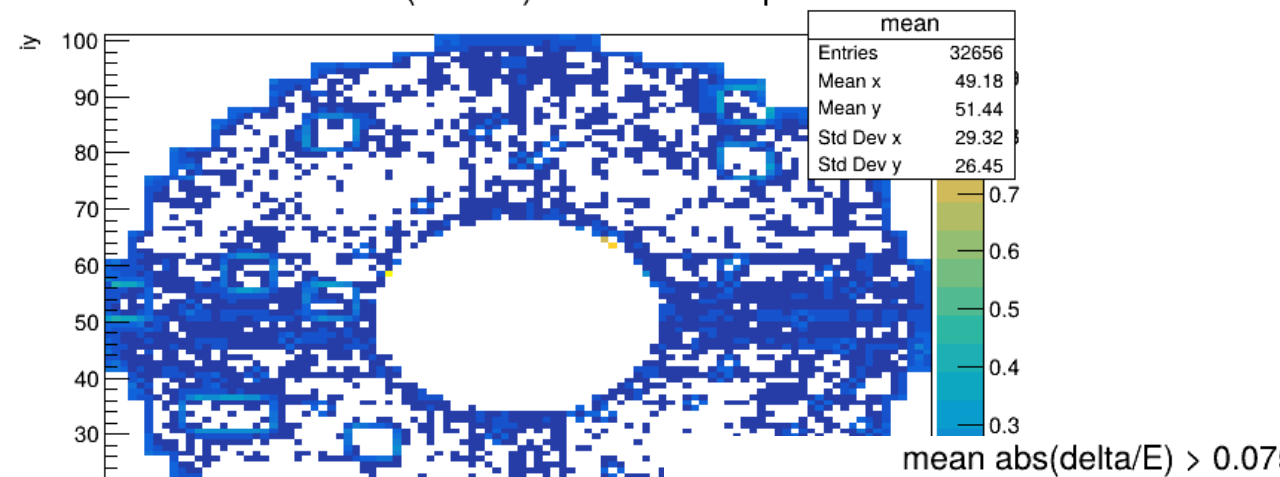
Error with and without shower variable + near-dead crystal cut offs; 2.5% removed



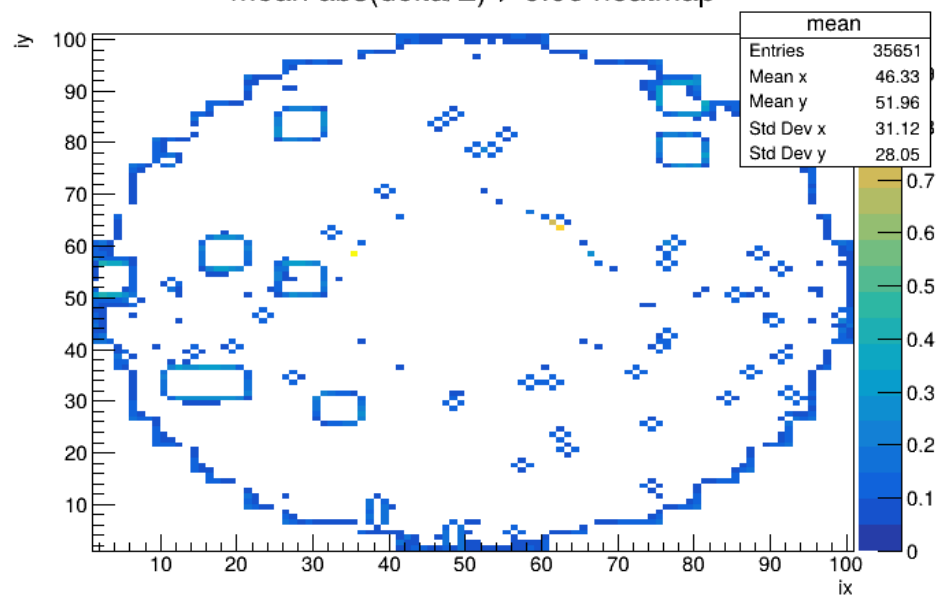
mean abs(delta/E) heatmap



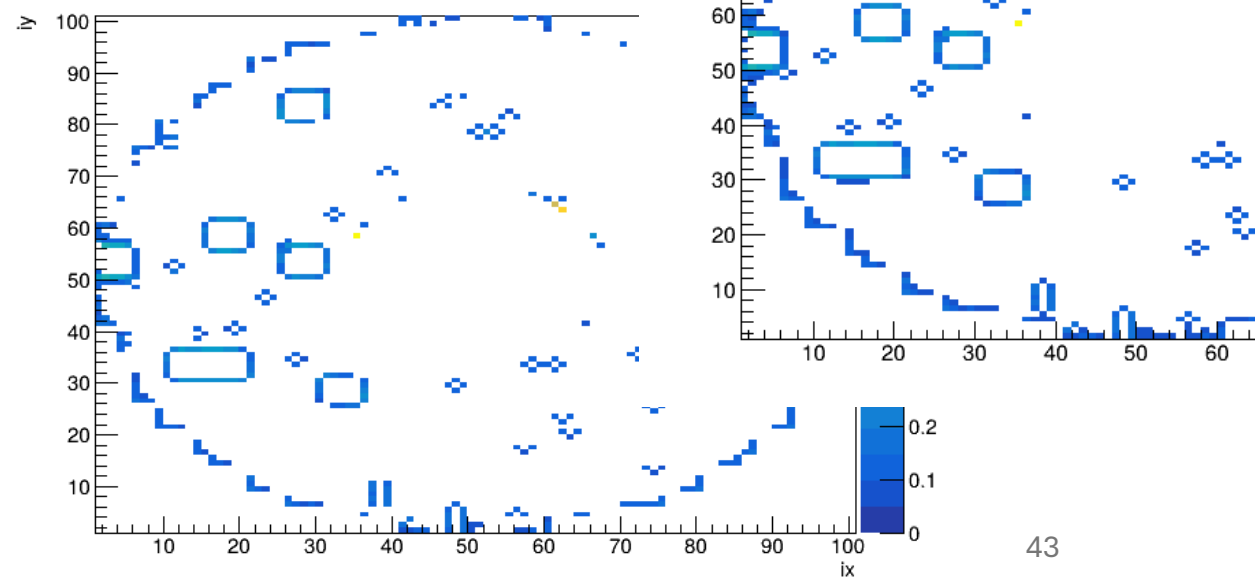
mean abs(delta/E) > 0.03 heatmap



mean abs(delta/E) > 0.06 heatmap



mean abs(delta/E) > 0.09 heatmap



BDT

- XGBoost

- Features:

- e1x5
 - e2nd
 - e2x5B, e2x5L, e2x5M, e2x5R, e2x5T
 - e5x5
 - eB, eL, eR, eT
 - eMax
 - eta ieta
 - iphi phi
 - pse
 - pt
 - r9
 - sieie

- Labels:

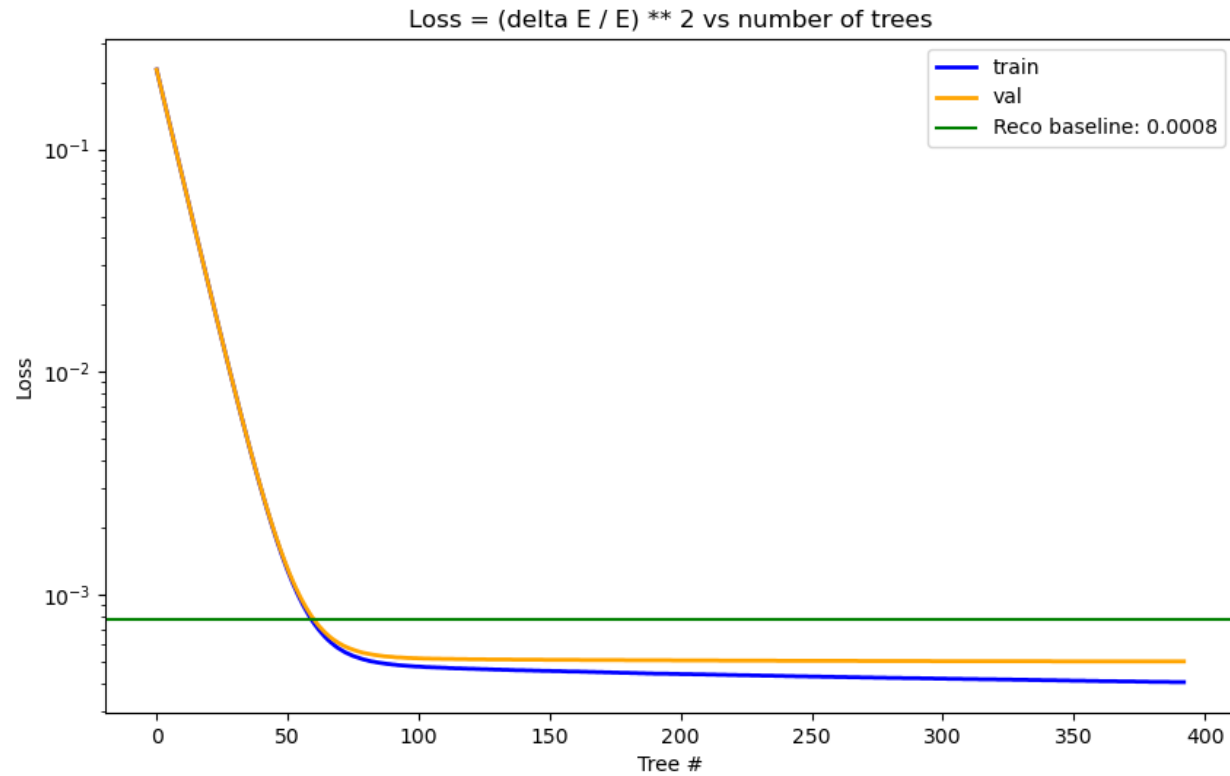
- $y_{\text{true}} = E_{\text{mc}} / E_{\text{reco}}$
 - $y_{\text{pred}} = \text{prediction on } y_{\text{true}}$

- Loss:

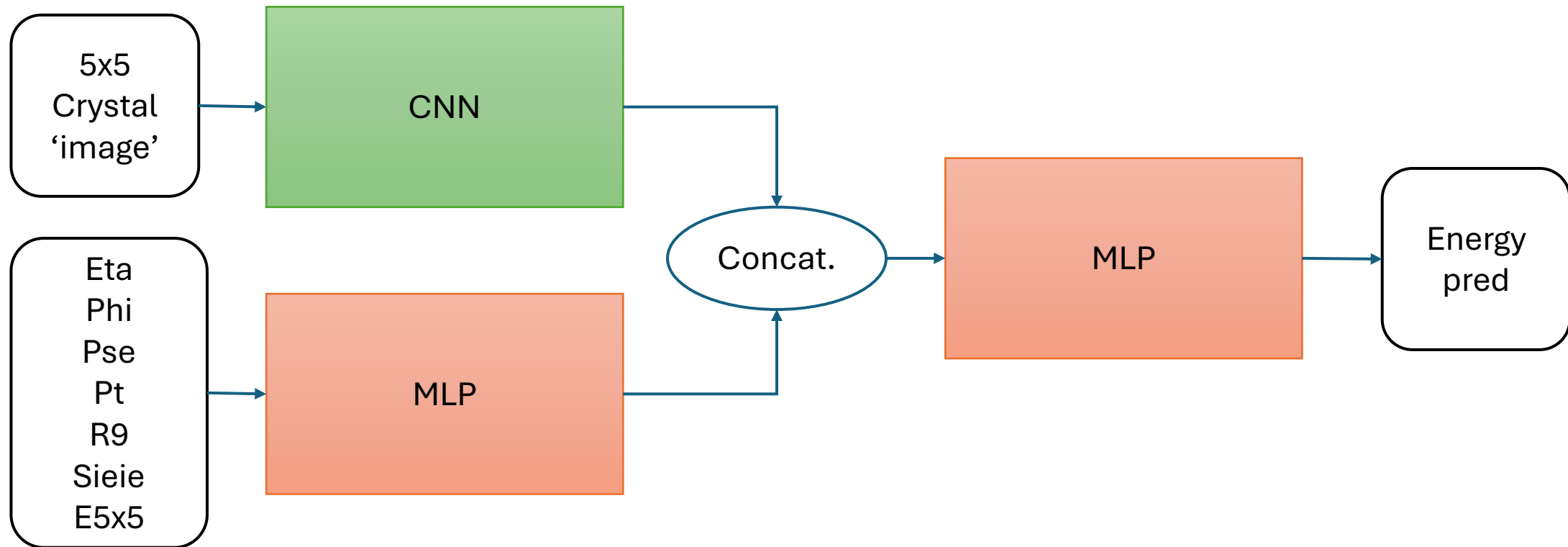
- $((E_{\text{pred}} - E_{\text{mc}}) / E_{\text{mc}})^2$
 - $= ((y_{\text{pred}} * E_{\text{reco}} - E_{\text{mc}}) / E_{\text{mc}})^2$
 - $= ((y_{\text{pred}} - E_{\text{mc}} / E_{\text{reco}}) / (E_{\text{mc}} / E_{\text{reco}}))^2$
 - $= ((y_{\text{pred}} - y_{\text{true}}) / y_{\text{true}})^2$

- No normalization (not necessary for BDT)

BDT (XGBoost) – Best model



Deep Learning: CNN – MLP Hybrid



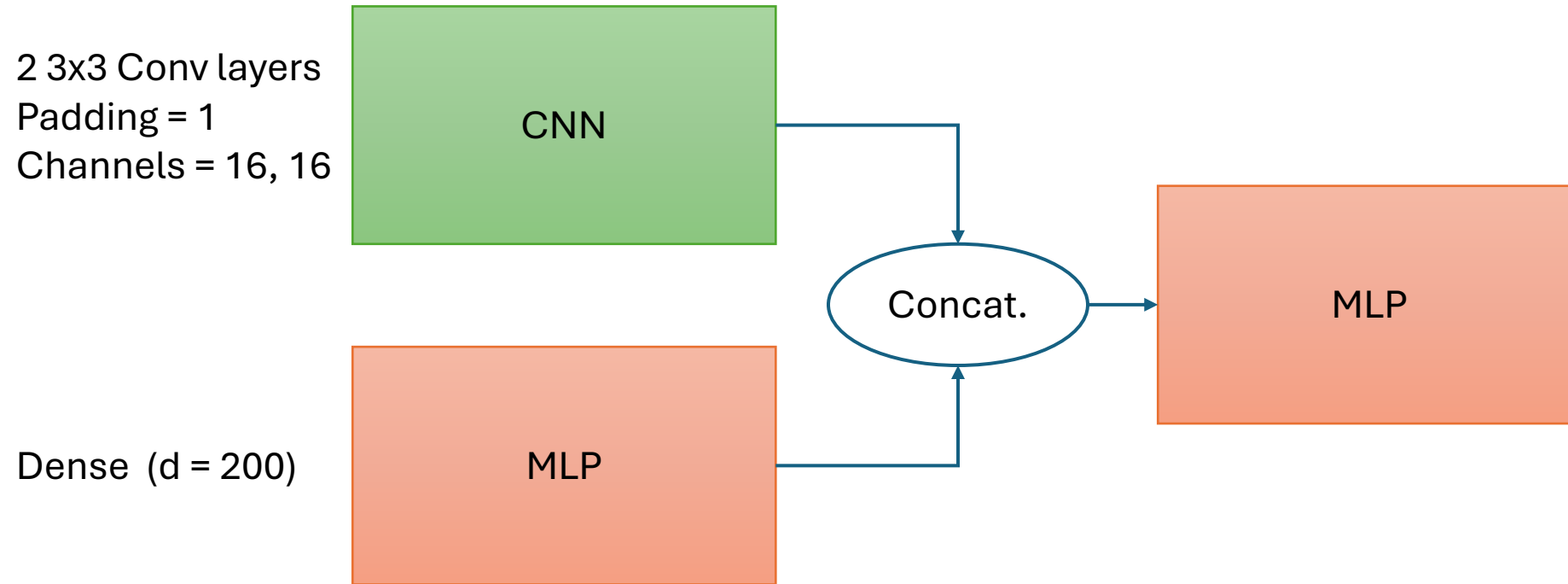
Attempted Normalizations

- Input image (X_{img}):
 - Normalise each image by its eMax
 - Log then normalise each image by its eMax
 - Sieie normalization: $(\max(0, 0.47 + \log(E_i/E_{5x5}))$
- Input features ($X_{tabular}$):
 - Mean std normalization
 - Transforming to $[0,1]$ (dividing by scale factor)
 - Taking log then applying one of the two above
 - Sin-cos embedding for phi
- Output labels (y):
 - Unscaled
 - Transforming to $[0,1]$ (dividing by scale factor)
 - Log
 - $Y = E_{mc} / E_{reco}$

Loss functions

- $L = ((E_{mc} - E_{reco} / E_{mc}))^{**2}$
- $L = (E_{mc} - E_{reco})^{**2}$
- $L = \log(\text{sig}^{**2}) + 0.5 * ((E_{mc} - \mu) / \text{sig})^{**2}$ (predicting μ and sig)
- First loss performed best

More on architecture



Train vs val for CNN

