

Direct optimization of the discovery significance for Machine Learning Analysis in CMS SUSY stop search

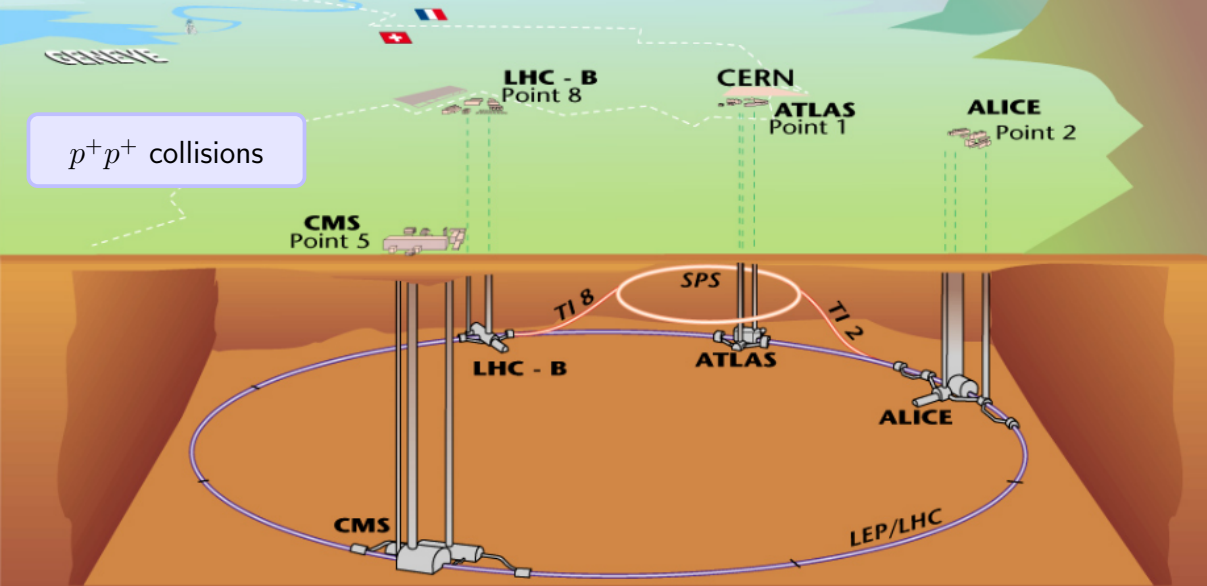
M.Shchedrolosiev

Content:

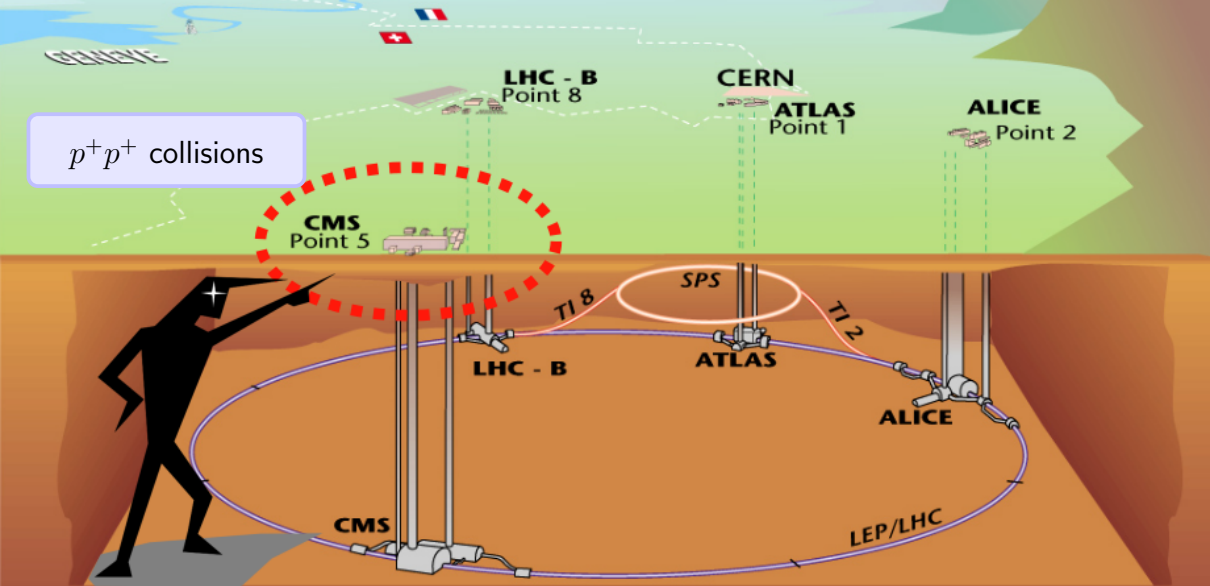
- CMS experiment
- Direct optimization of the discovery significance
- XGBoost (Boosted decision tree approach)
- Stop SUSY model Monte Carlo
- Results



Supervised by: A.Elwood, D.Krücker, I.Melzer-Pellmann, O.Turkot
Thanks for contribution to: C.Contreras



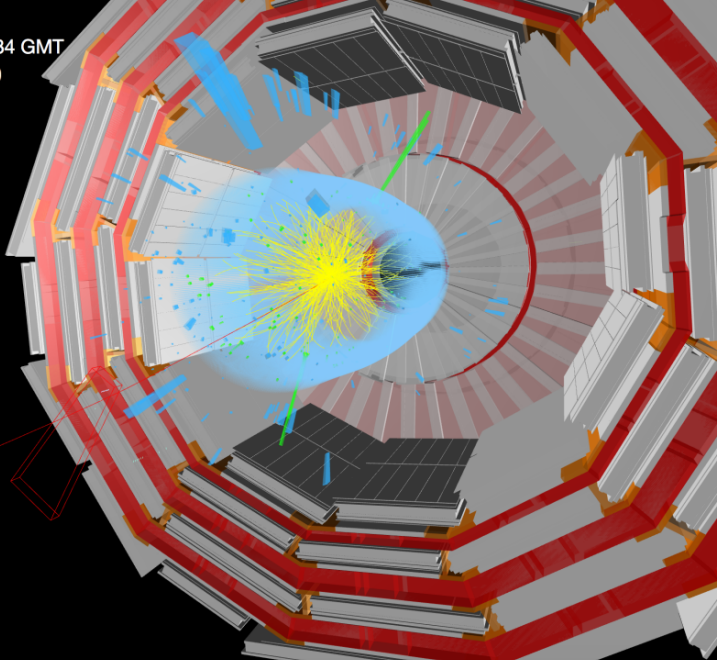
p^+p^+ collisions





Purposes:

- SUSY events with 1 lepton and multiple jets in pp collisions at $\sqrt{s} = 13$ TeV
- Train algorithm that evaluate SUSY events from all DATA on Monte Carlo
- Use this algorithm for CMS data

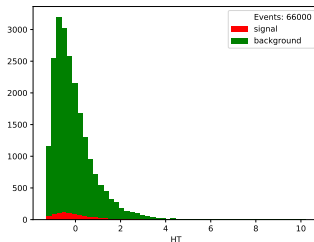
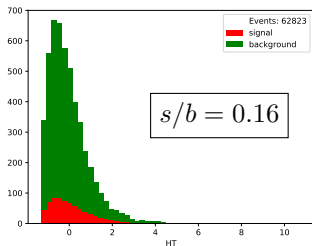


Direct optimization of the discovery significance

- When searching for new physics the most important is the significance of signal counts over background counts, but purity of the background classification is not very important

Standard ML approach

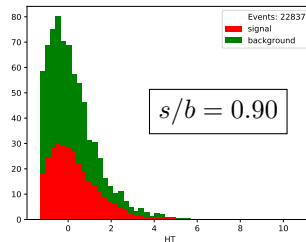
Signal prediction



input data

with a HEP approach

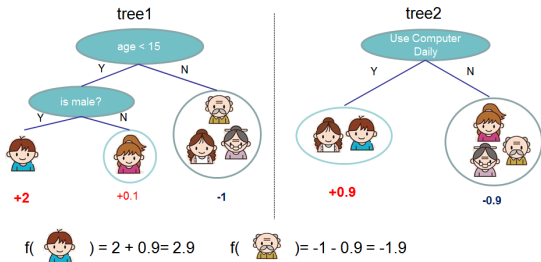
Signal prediction



- To train Machine Learning algorithms at HEP approach we want to directly maximize discovery significance, not accuracy or ROC curve area, to get a sample where the signal dominates in signal prediction

XGBoost (Boosted decision tree approach)

Example: does person play PC games?



- Prediction is sum of scores predicted by each of the tree

- Gradient Tree Boosting :

- $\hat{y}_i^{(t)}$ is prediction of the i -th instance at the t -th iteration, Ω - penalizes the complexity of the model
- XGBoost uses second order approximation
- Additive Training (Boosting)
- Minimize every next tree f_t

$$\mathcal{L}^{(t)} \simeq \sum_{i=1}^n [l(y_i, \hat{y}^{(t-1)}) + g_i f_t(x_i)] + \Omega(f_t)$$

where $g_i = \partial_{\hat{y}^{(t-1)}} l(y_i, \hat{y}^{(t-1)})$ are the gradient statistics on the loss function.

Direct optimization of the discovery significance¹

- The standard approach is to maximize accuracy through minimizing Binary cross entropy:

$$C = -\frac{1}{n} \sum_x [y^{true} \ln y^{pred} + (1 - y^{true}) \ln(1 - y^{pred})]$$

- equivalent to minimize logarithmic likelihood for binomial model

- To train neural networks in HEP approach, one can design a loss function based around the direct optimization of the discovery significance, for instance maximize Asimov discovery significance (minimize loss $l_{Asimov} = 1/Z_A$):

$$Z_A = \sqrt{2((s+b) \ln[\frac{(s+b)(b+\sigma_b^2)}{b^2 + (s+b)\sigma_b^2}] - \frac{b^2}{\sigma_b^2} \ln[1 + \frac{\sigma_b^2 s}{b(b+\sigma_b^2)})] \rightarrow \frac{s}{\sqrt{s+b}}}$$

- **s** - correctly classified signal events
- **b** - incorrectly classified background events
- σ - systematic uncertainty

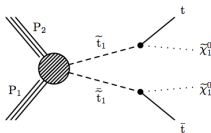
$$s = W_s \sum_i^{N_{batch}} y_i^{pred} \times y_i^{true},$$

$$b = W_b \sum_i^{N_{batch}} y_i^{pred} \times (1 - y_i^{true}),$$

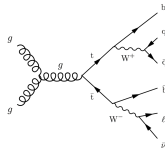
¹arXiv:1007.1727v3

Stop SUSY model MC²

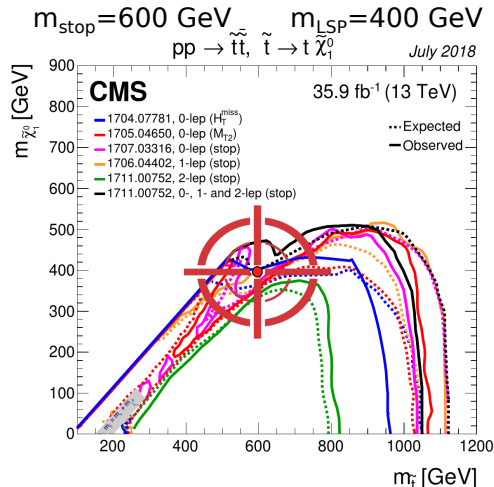
- To test this approach look at the stop SUSY model that is close to the edge of exclusion at 30fb^{-1} of 13 TeV LHC data
- Consider top/anti-top quark pairs as the background
- Taken 1M events of signal and background with Pythia and Delphes with basic selection criteria: 1 lepton $p_T \geq 40$ GeV, 4 jets $p_T \geq 30$ GeV, at least 1 b-tagged jet



Signal

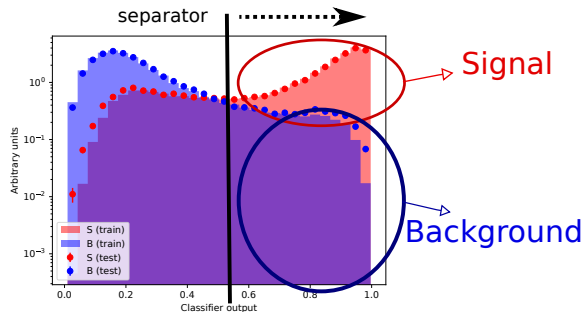


Background



²<https://indico.desy.de/indico/event/21116/contribution/0/material/slides/0.pdf>

XGBoost (Classifier output)



- Take events that were classified with value more than some fixed value of separator
- Calculate how much signal and background there
- Calculate value of functional that we want to optimize (for instance: Asimov significance)

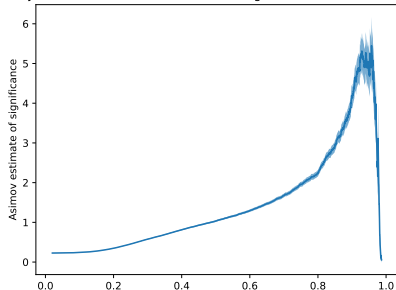
$$Z_A = \sqrt{2((s+b) \ln[\frac{(s+b)(b+\sigma_b^2)}{b^2+(s+b)\sigma_b^2}] - \frac{b^2}{\sigma_b^2} \ln[1 + \frac{\sigma_b^2 s}{b(b+\sigma_b^2)}])}$$

Note that:

3σ - observation

5σ - discovery

Systematic 0.2, s: 56.4, b:26.3, best significance is 5.45 +/- 0.72



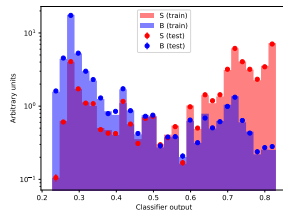
XGBoost out-of-box (Asimov estimate scorer, Z_A)

- Much worse comparing to the neural network²

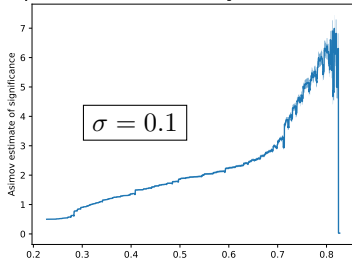
σ	0.1	0.3	0.5
Asimov NN	10.7	6.8	4.8
Asimov XGBoost	7.0	3.9	2.6

² arXiv:1806.00322

- Has to be tuned basing on Asimov significance
- Early stopping is used basing on Asimov scorer

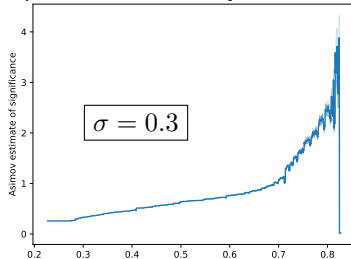


Systematic 0.1, s: 66.2, b:39.0, best significance is 6.99 +/- 0.45



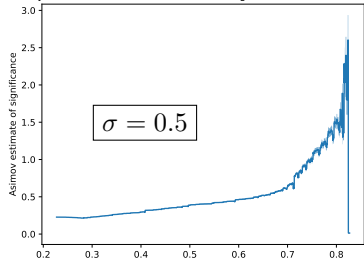
M. Shchedrolosiev

Systematic 0.3, s: 37.5, b:17.6, best significance is 3.88 +/- 0.45



CMS DESY GROUP

Systematic 0.5, s: 37.5, b:17.6, best significance is 2.6 +/- 0.34

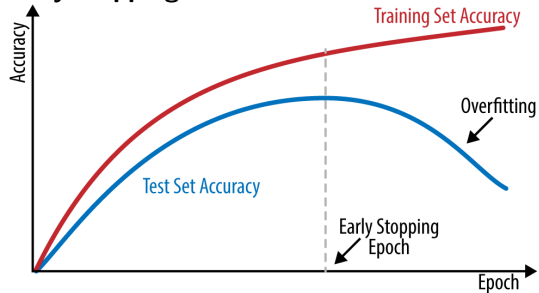


September 6, 2018

9 / 20

XGBoost (tuning and implementation of objective loss function)

Early stopping:



- Asimov score as a metrics for early stopping procedure
- To avoid an error in early stopping continuously differentiable function of Asimov function was used

XGBoost objective function

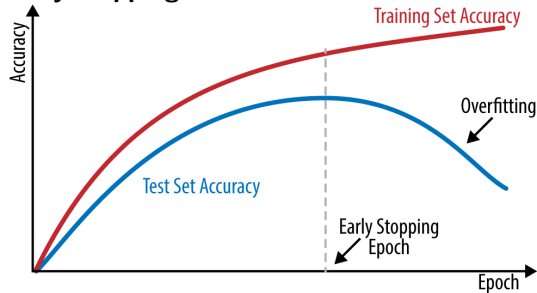
$$\mathcal{L}^{(t)} \simeq \sum_{i=1}^n [l(y_i, \hat{y}^{(t-1)}) + g_i f_t(i)] + \Omega(f_t)$$

where we define loss function as: $l_{Asimov} = 1/Z_A$ and gradient:

$$g_i = \frac{\partial l(y_i, \hat{y}^{(t-1)})}{\partial \hat{y}^{(t-1)}}$$

XGBoost (tuning and implementation of objective loss function)

Early stopping:



- Asimov score as a metrics for early stopping procedure
- To avoid an error in early stopping continuously differentiable function of Asimov function was used

XGBoost objective function

$$\mathcal{L}^{(t)} \simeq \sum_{i=1}^n [l(y_i, \hat{y}^{(t-1)}) + g_i f_t(i)] + \Omega(f_t)$$

where we define loss function as: $l_{Asimov} = 1/Z_A$
and gradient:

$$g_i = \frac{\partial l(y_i, \hat{y}^{(t-1)})}{\partial \hat{y}^{(t-1)}}$$

$$g_i = -Z_A^{-2} \left(\frac{\partial Z_A}{\partial s} W_s y_i + \frac{\partial Z_A}{\partial b} W_b (1 - y_i) \right)$$

XGBoost (tuning and implementation of objective loss function)

**Binary cross entropy
as a minimization function in fit**

$$C = -\frac{1}{n} \sum_x [y \ln a + (1 - y) \ln(1 - a)]$$

**Asimov significance
as a minimization function in fit**

$$Z_A = [2((s + b) \ln[\frac{(s+b)(b+\sigma_b^2)}{b^2+(s+b)\sigma_b^2}] - \frac{b^2}{\sigma_b^2} \ln[1 + \frac{\sigma_b^2 s}{b(b+\sigma_b^2)}])]^{1/2}$$

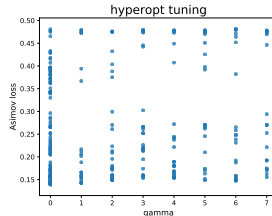
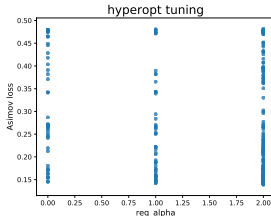
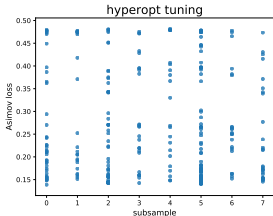
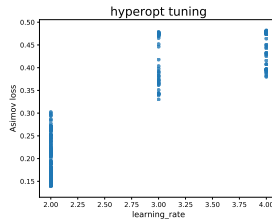
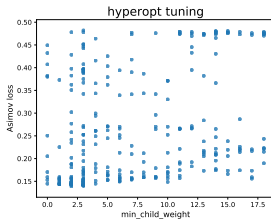
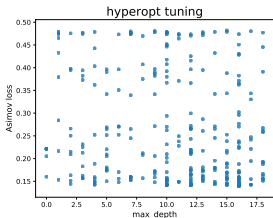
Tuning

- Changed and adjust separation facet for predicting of class
- Used Asimov score as a metrics for early stopping approach
- Used Asimov score as a metrics for a hyperparameter tuning

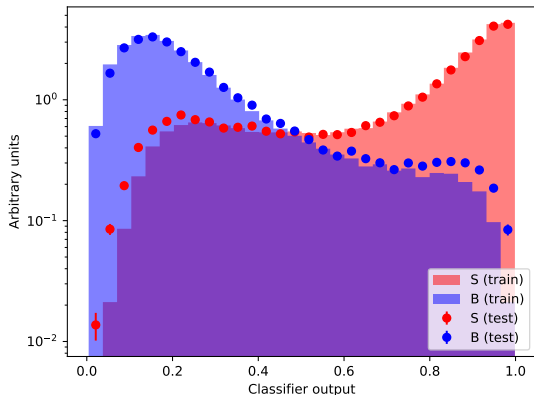
Goal is to compare both after tuning!

XGBoost (Hyperparameter tuning)

- Put Asimov loss function as an evaluation function for hyper-parameter tuning $l_{Asimov} = 1/Z_A$. Find optimal parameters for XGBoost Classifier:

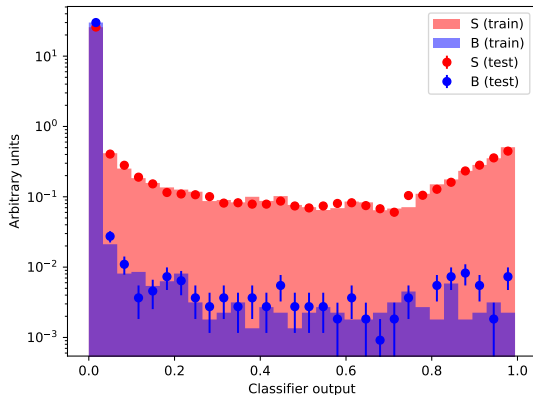


XGBoost (Train and Test set comparison)



Binary cross entropy loss function

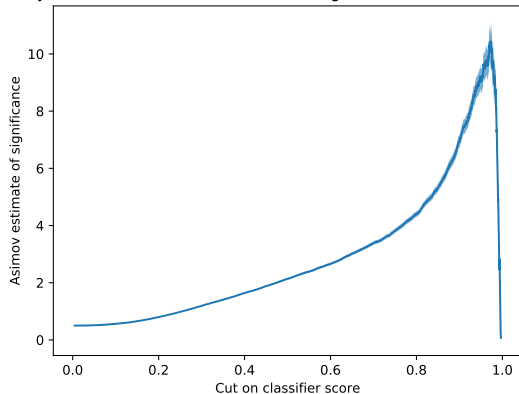
Asimov significance demonstrates a good separation of signal and background in a wide range of probability values



Asimov loss function

XGBoost ($\sigma = 0.1$)

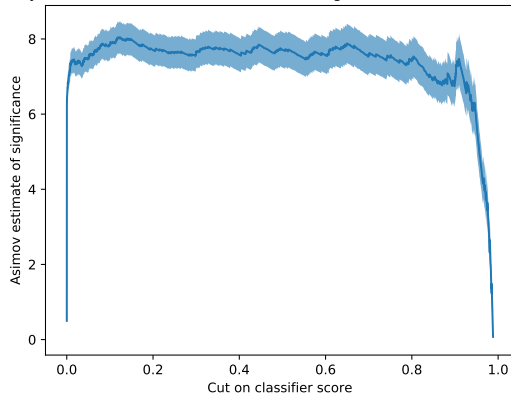
Systematic 0.1, s: 88.2, b:26.4, best significance is 10.44 +/- 0.68



Binary cross entropy loss function

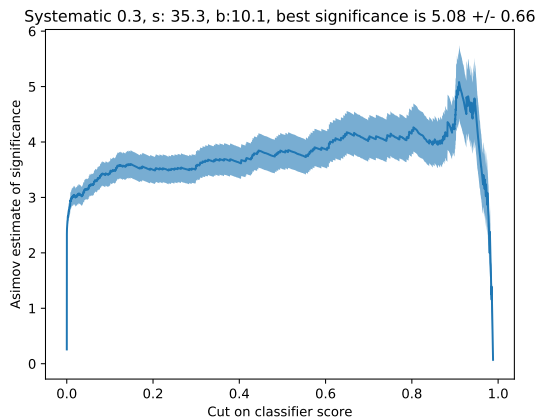
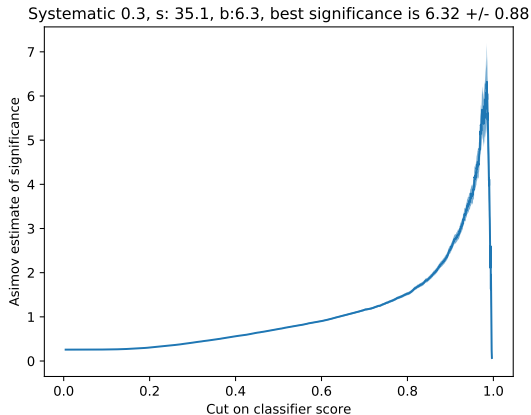
Binary cross entropy demonstrates better result than Asimov significance as an objective function for small uncertainty

Systematic 0.1, s: 121.5, b:76.1, best significance is 8.05 +/- 0.42



Asimov loss function

XGBoost ($\sigma = 0.3$)



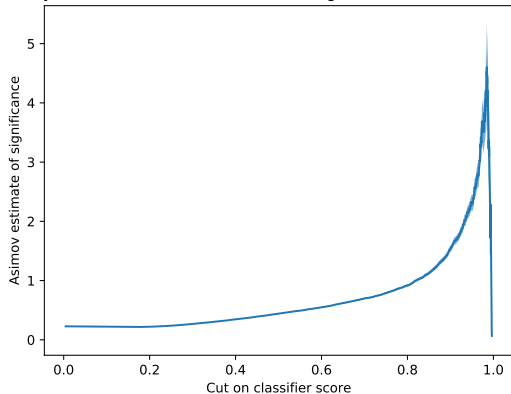
Binary cross entropy loss function

The difference between Binary cross entropy and Asimov significance diminishes with the increase of uncertainty

Asimov loss function

XGBoost ($\sigma = 0.5$)

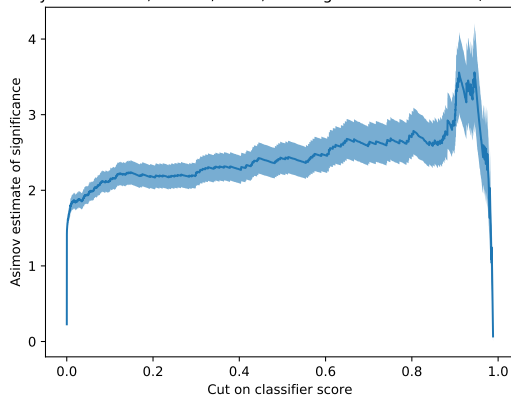
Systematic 0.5, s: 35.1, b:6.3, best significance is 4.6 +/- 0.77



Binary cross entropy loss function

Binary cross entropy and Asimov loss are comparable within uncertainty ranges for $\sigma = 0.5$

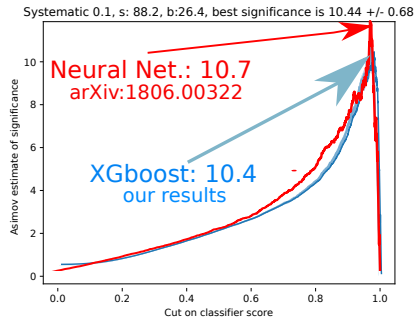
Systematic 0.5, s: 22.1, b:5.7, best significance is 3.56 +/- 0.65



Asimov loss function

Summary

- Implemented XGBoost classifier in framework for the SUSY analysis
<https://github.com/shedprog/hepML>
- Implemented Asimov loss as an objective function and evaluating metrics for an early stopping and for hyperparameter tuning
- Showed that tuned XGBoost with the default binary cross entropy performs almost as good as the neural network
- Showed that using of Asimov loss as an objective function lead to small improvement (but very close to binary cross entropy for high systematic)





THANK
YOU!