

AVANTIX

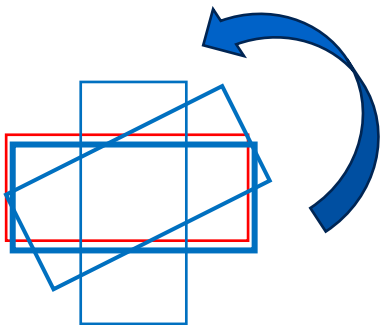
Avancement Simulateur IQ

AVANTIX

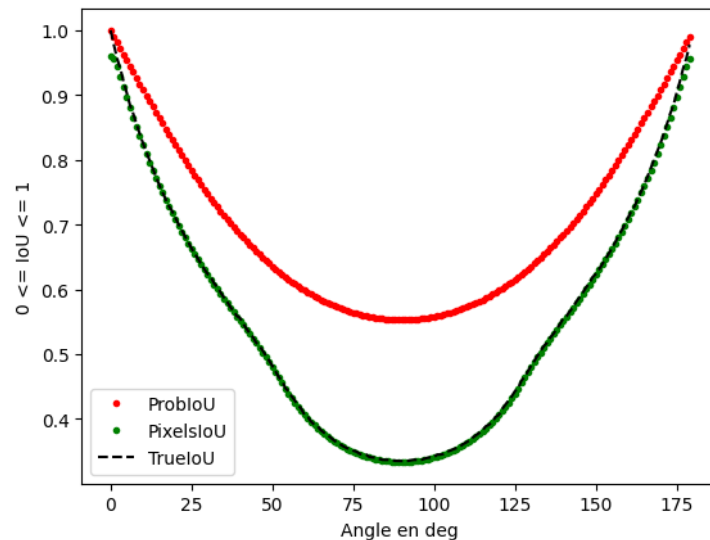
ProbloU vs PixelsloU

ProbloU vs PixelsloU

Comparaison entre les deux par rapport à la vrai IoU



On fait tourner un rectangle autour de son centre et on calcule les 3 métriques entre le rouge et le bleu qui tourne



PixelsloU approche très bien la vrai IoU

ProbloU vs PixelsloU

En lançant un apprentissage avec PixelsloU par substitution du ProbloU la fonction de coût est constante durant l'apprentissage.

Il faudrait utiliser la fonction de coût qui accompagne la PixelsloU et qui est détaillée dans le papier → **Objectif non prioritaire**

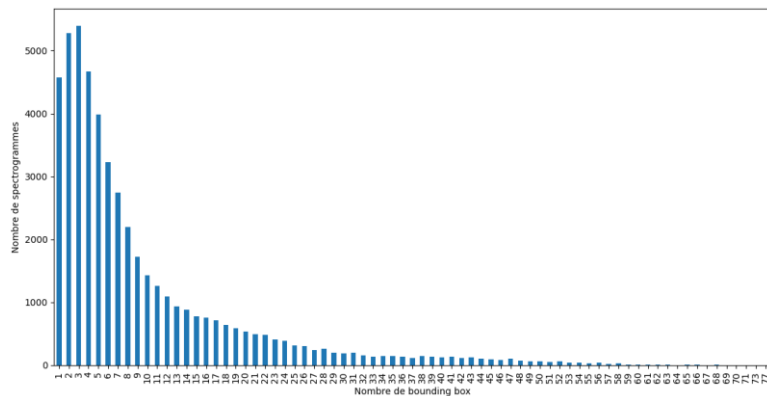
Utilisation de ProbloU avec cap des rapports d'aspects

AVANTIX

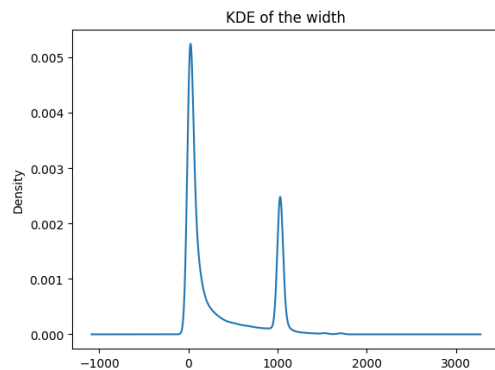
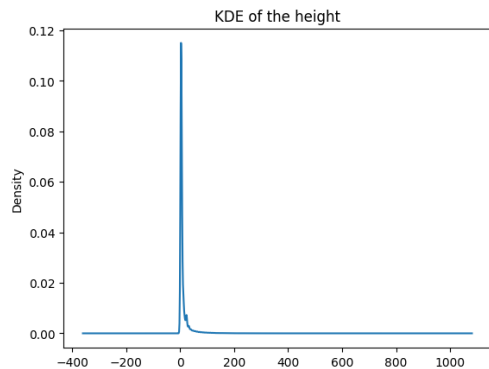
BT300

- 49 529 images 80% train, 20% val et 20% test
- $\sigma_{\text{noise}} = 1\text{mV}$, $f_{\text{c_rec}} = 2,5\text{GHz}$
- RADAR & RADAR+COM : 50 / 50
- $f_{\text{s_rec}} = [4, 20, 100, 500, 1\text{e}3, 4\text{e}3]\text{MSps}$
- $n_{\text{perseg}} = [128, 256, 512, 1024]$, $n_{\text{fft}} = 1024$, $n_{\text{overlap}} = n_{\text{perseg}}/2$
- $\text{snr_eff_db_min} = -5\text{dB}$
- Mise à l'échelle : $\text{min} = 1^{\text{er}} \text{ décile}$, $\text{max} = 3000 * 1^{\text{er}} \text{ décile}$

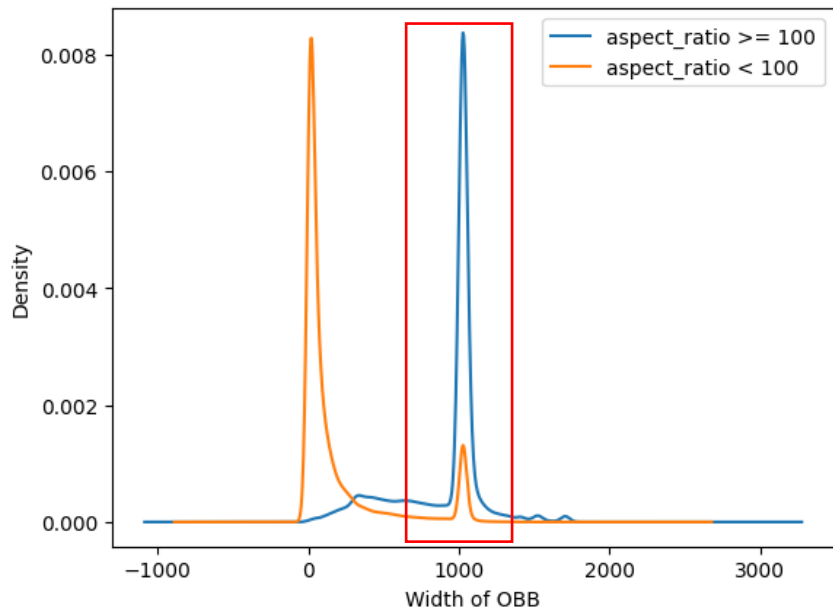
BT300



Présence de $h=0$



Cap du rapport d'aspect à 100

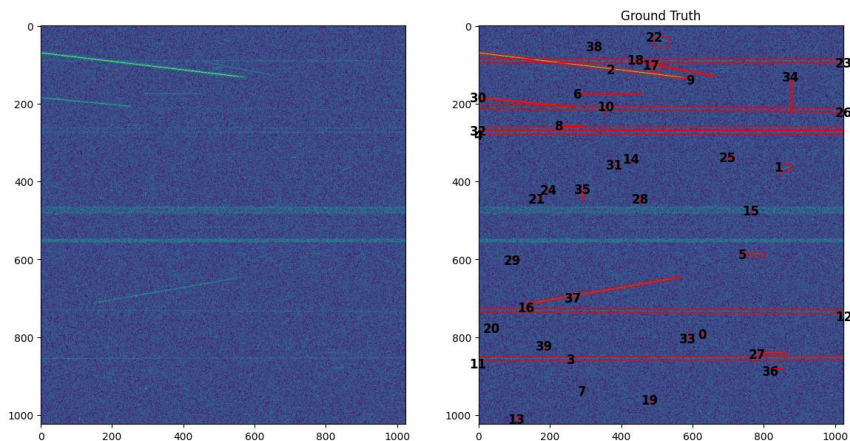


On remarque que :

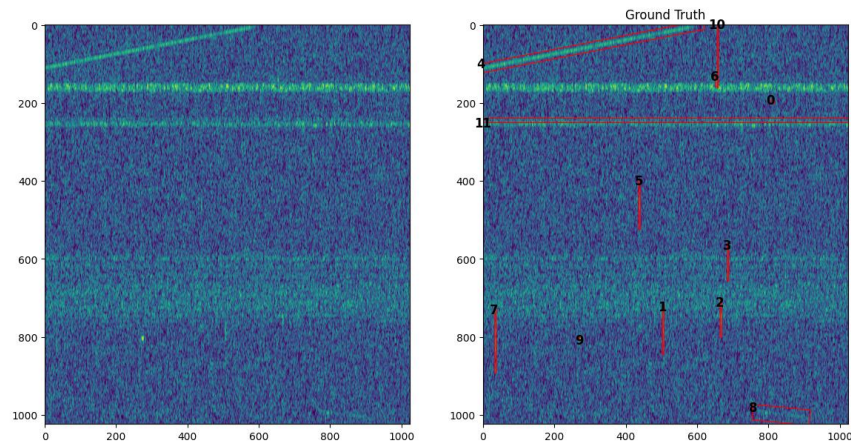
- Quasiment toutes les obboxes ayant un $AR > 100$ ont des largeurs égales à la largeur de l'image
- Il existe tout de même des obboxes ayant des largeurs égales à la largeur de l'image qui ont un rapport d'aspect inférieur

Cap du rapport d'aspect à 100

rec_17795_fs_4000.0MSps_nperseg_1024.txt



rec_31756_fs_500.0MSps_nperseg_128.txt



Cap du rapport d'aspect à 50 → Présence de NaN mais moins fréquente

Après investigation :

- $ab-c^2$ peut être négatif sans que $t3_exp$ le soit
- $t3_exp$ peut être négatif car $(ab-c^2)_{gt} < 0$
- $t3_exp$ peut être négatif sans que $ab-c^2 < 0$
- Lorsque $ab-c^2 < 0$, ça n'est que pour des obb de la GT
- Jamais vu $a < 0$ ou $b < 0$



Solution : Cap du $t3_exp$ à 10^{-7}

Et en regardant une image provoquant un NaN, je n'ai rien vu d'inhabituelle/aberrant/problématique

Rappel

$$t3_exp = \frac{(a_1 + a_2)(b_1 + b_2) - (c_1 + c_2)^2}{4\sqrt{(a_1b_1 - c_1^2)(a_2b_2 - c_2^2)}}$$

AVANTIX

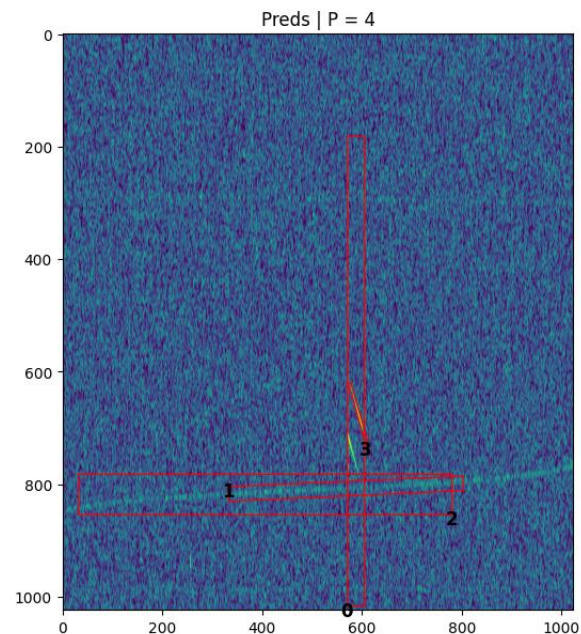
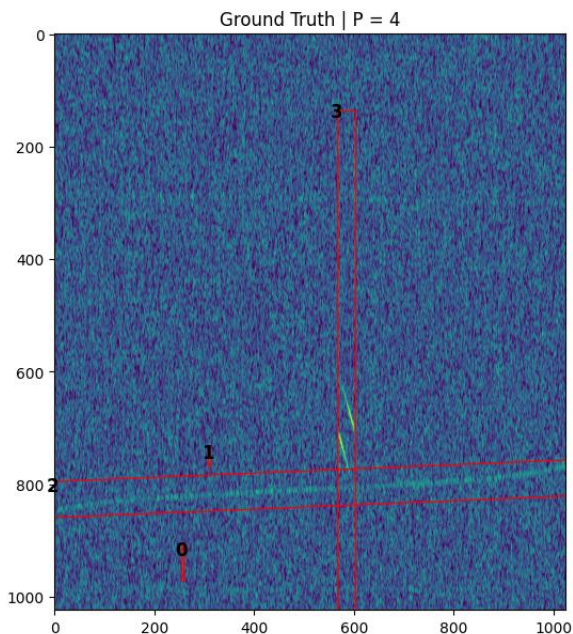
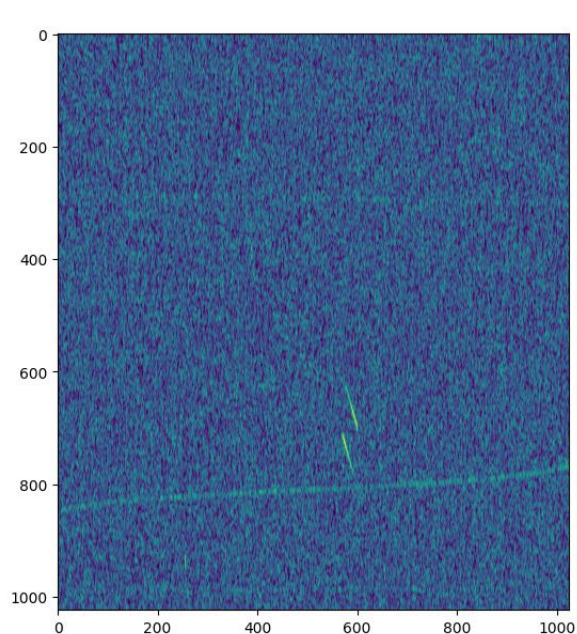
Meilleur modèle du moment

Détail du meilleur modèle

- Fine tuning des HP qui a débuté le 04/06 à 10h38
- Choisis parmi 5 modèles entraînés uniquement (~5h de calcul pour 1 modèle !)
- Nano
- Sans cap du t3_exp
- Val_Box_loss à 1,3117
- Entraîné sur 30% de BT300 (~18k exemples en apprentissage)

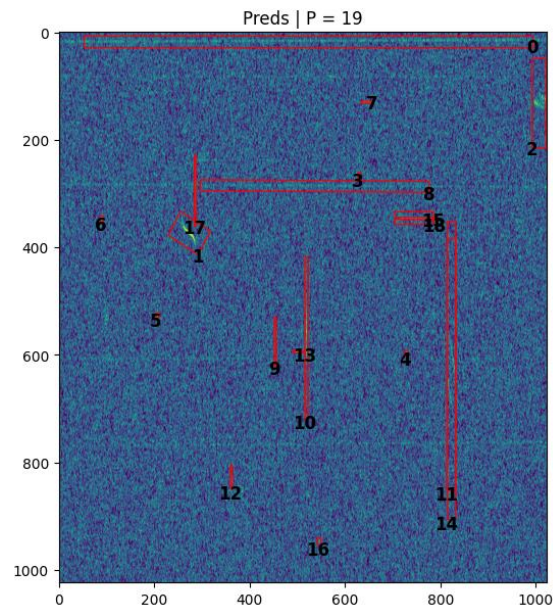
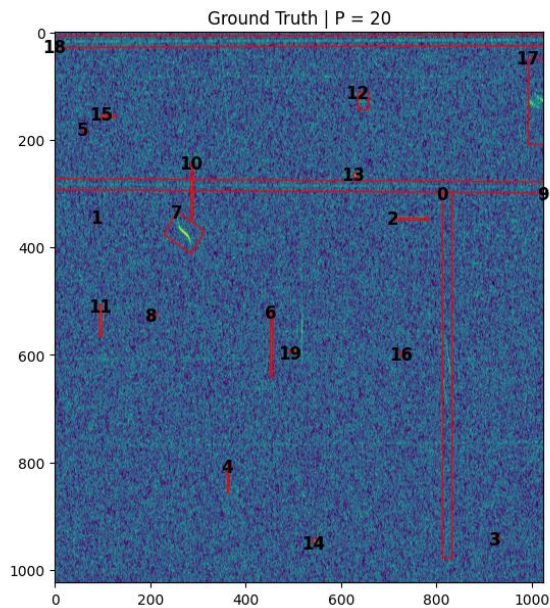
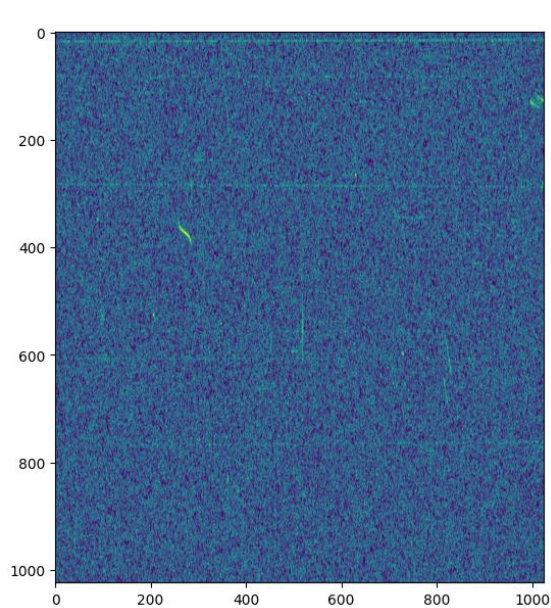
Inférence sur test

seed:0, rec_8487_fs_1000.0MSps_nperseg_128.txt, val_loss=1.4407



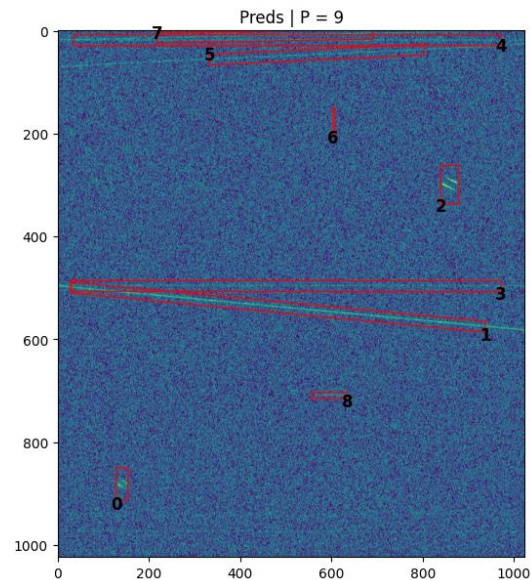
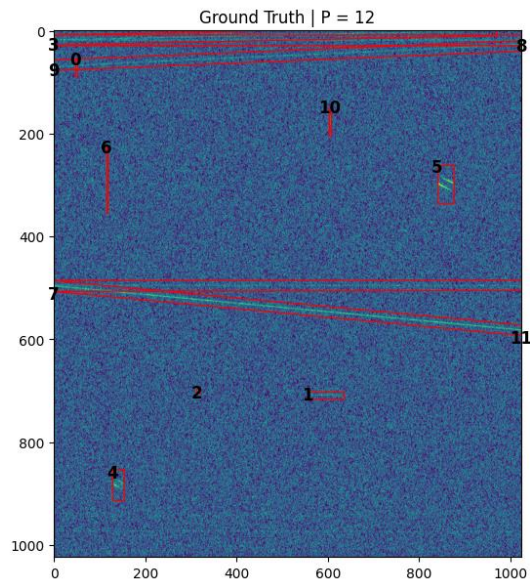
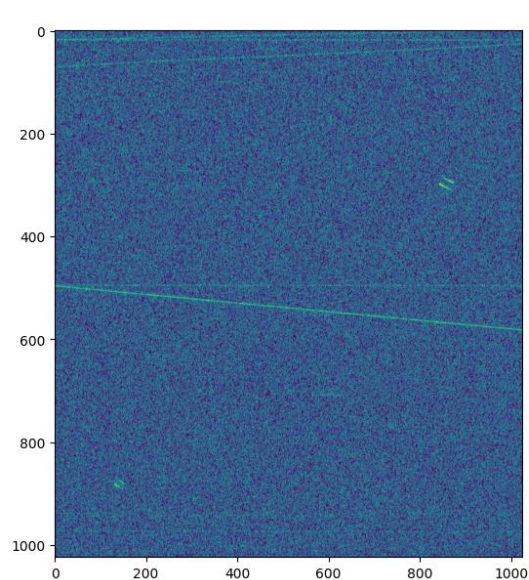
Inférence sur test

seed:1, rec_50342_fs_1000.0MSps_nperseg_256.txt, val_loss=1.4407



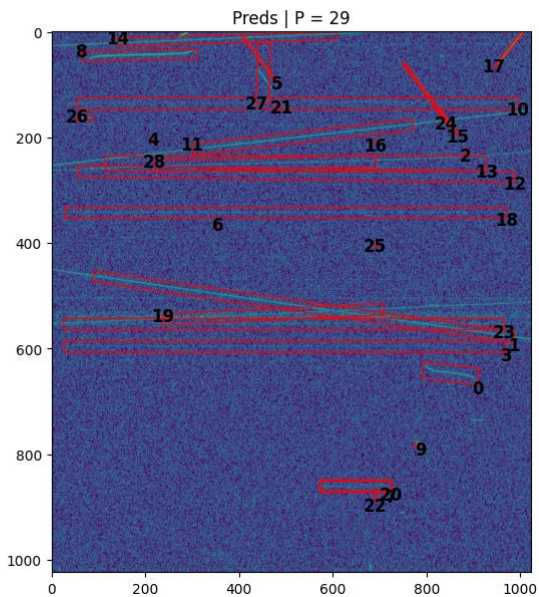
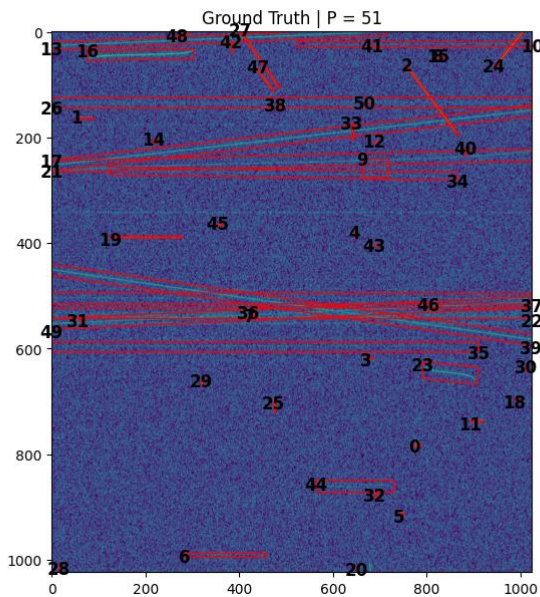
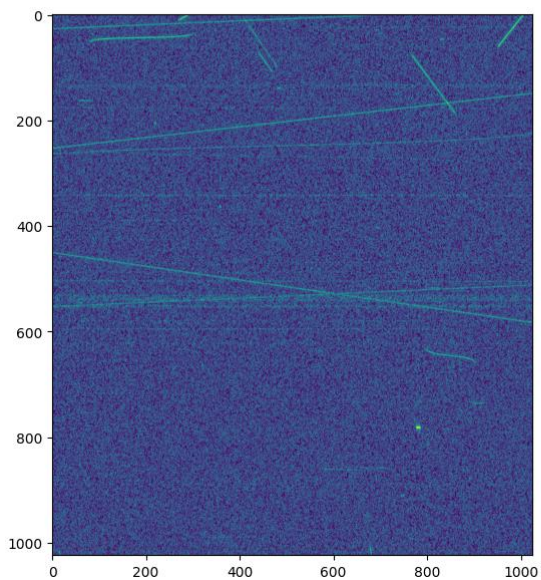
Inférence sur test

seed:2, rec_8044_fs_1000.0MSps_nperseg_512.txt, val_loss=1.4407



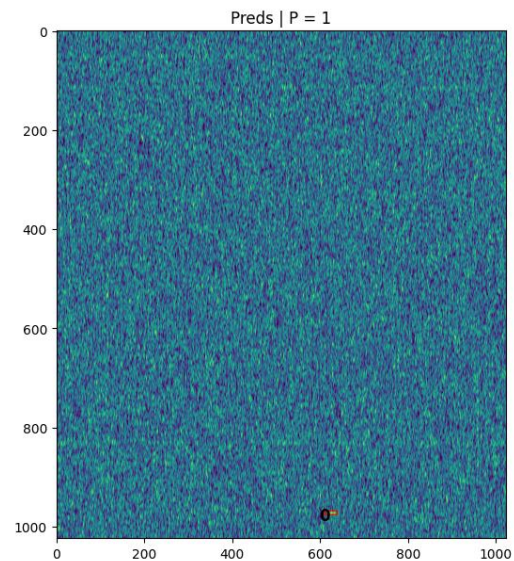
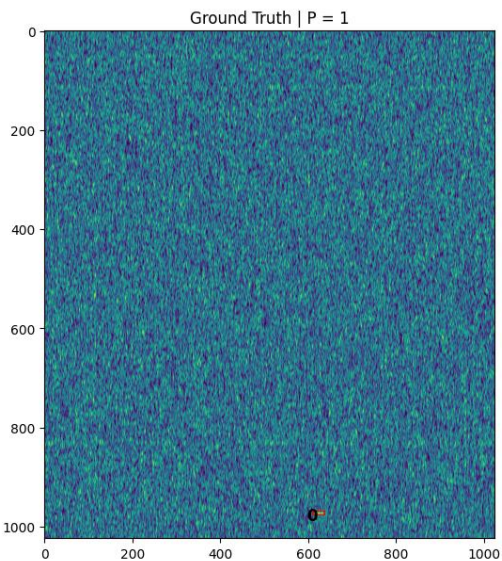
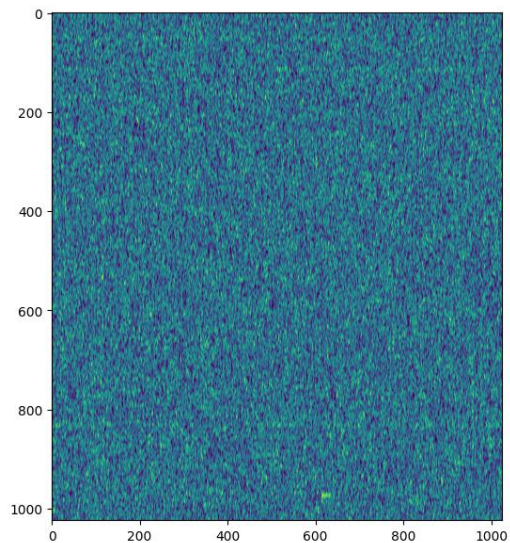
Inférence sur test

seed:3, rec_4125_fs_4000.0MSps_nperseg_512.txt



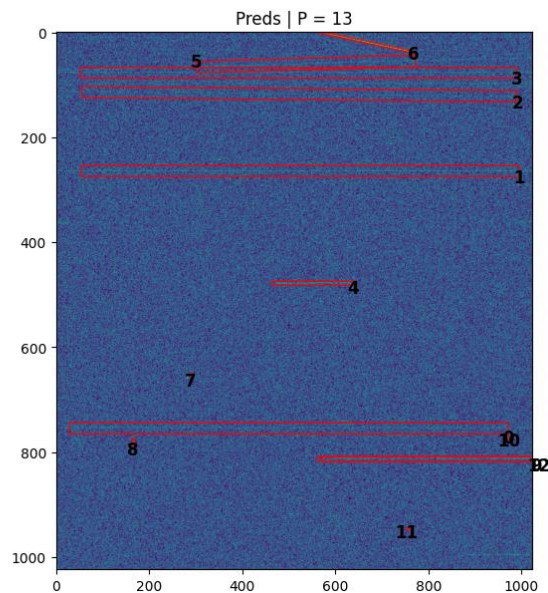
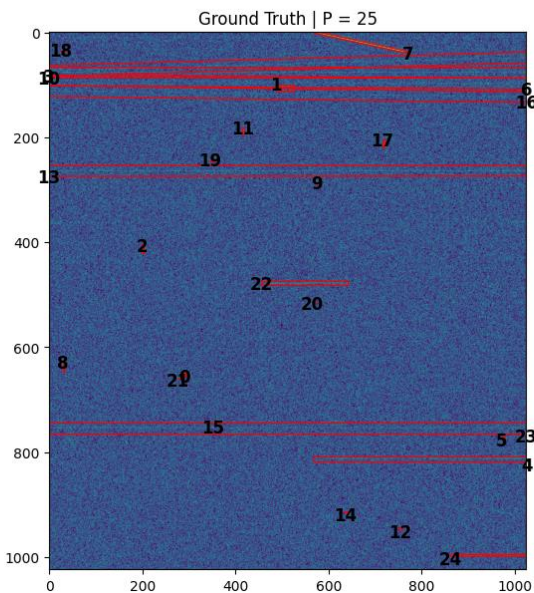
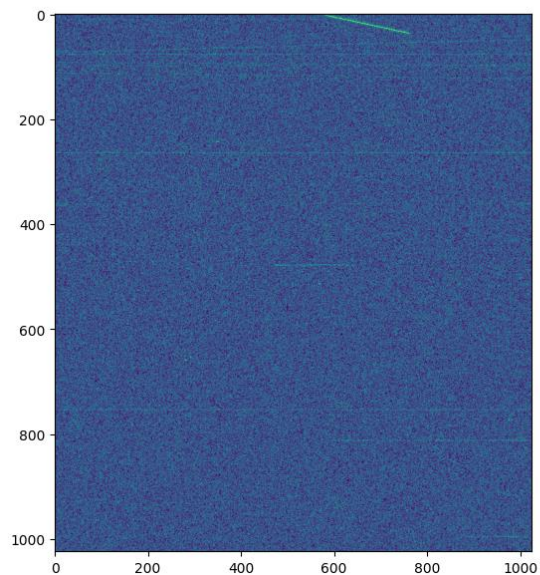
Inférence sur test

seed:4, rec_33981_fs_4000.0MSps_nperseg_128.txt



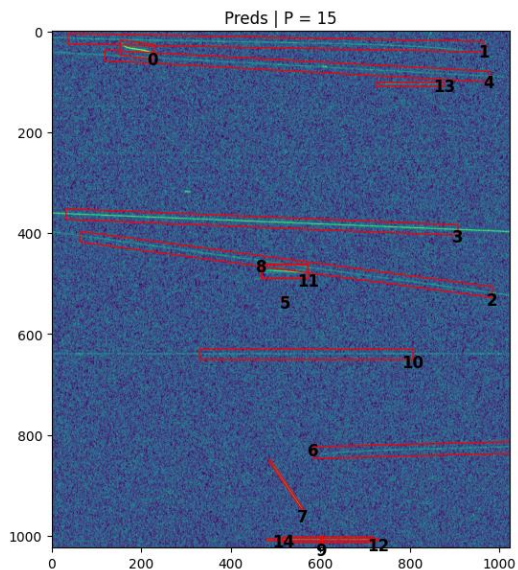
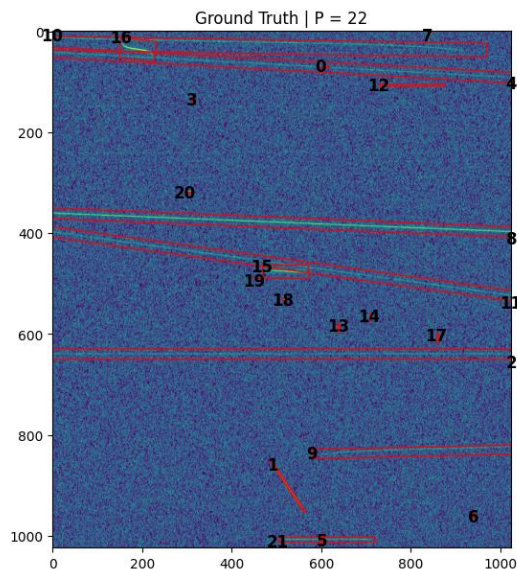
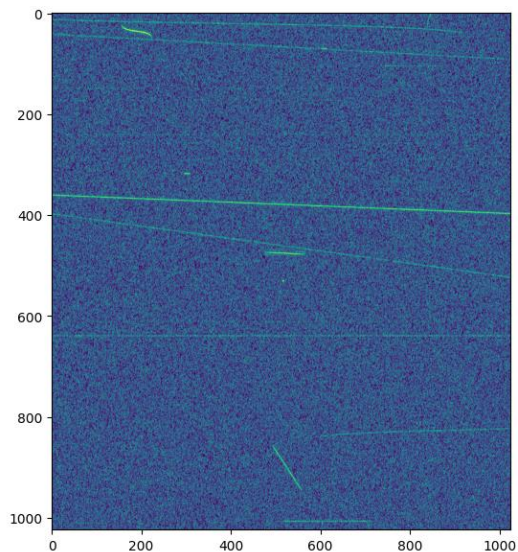
Inférence sur test

seed:5, rec_179_fs_4000.0MSps_nperseg_1024.txt



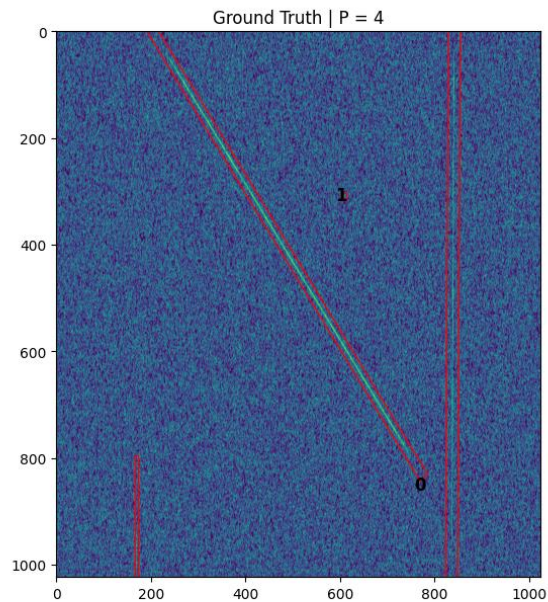
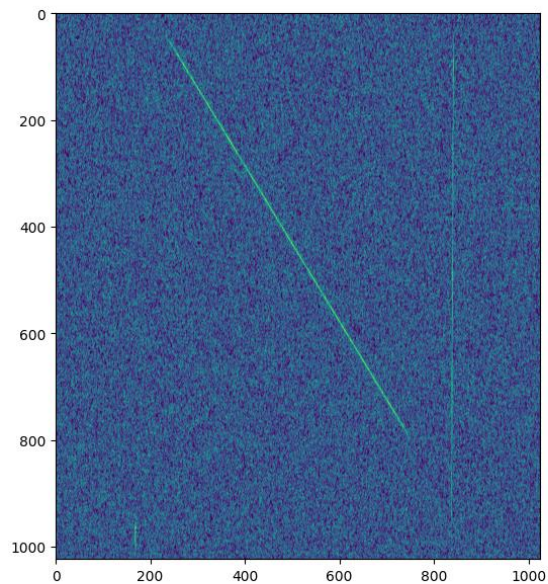
Inférence sur test

seed:6, rec_31904_fs_4000.0MSps_nperseg_512.txt



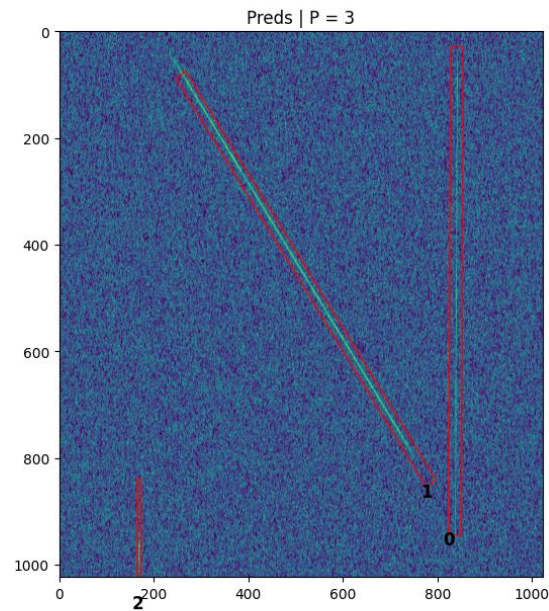
Inférence sur test

seed:7, rec_46199_fs_4.0MSps_nperseg_256.txt



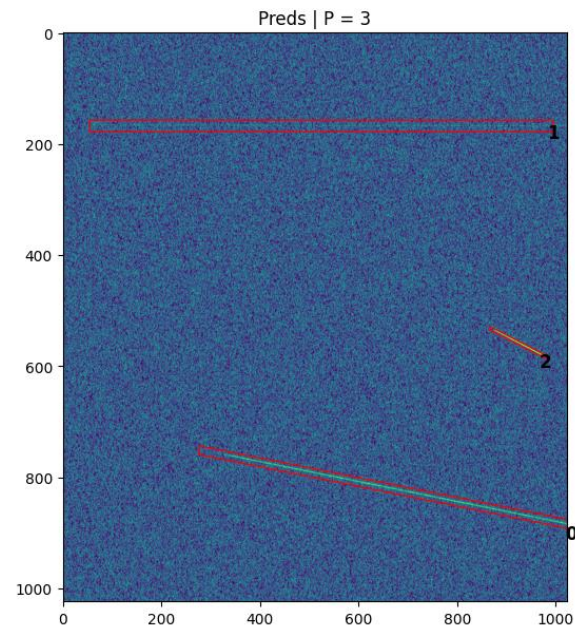
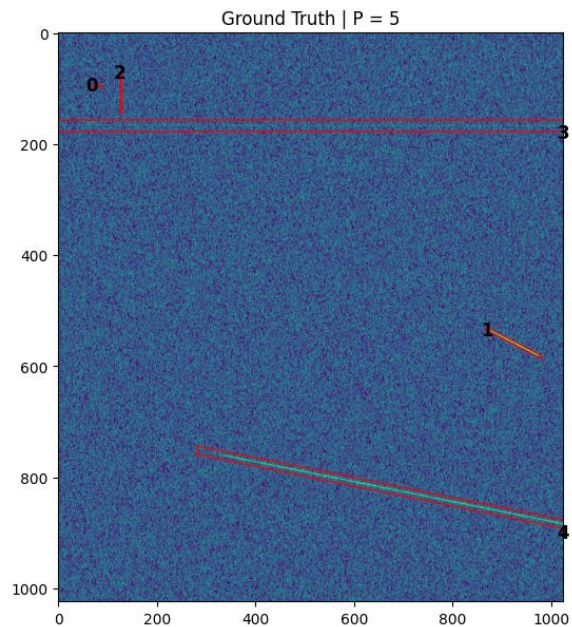
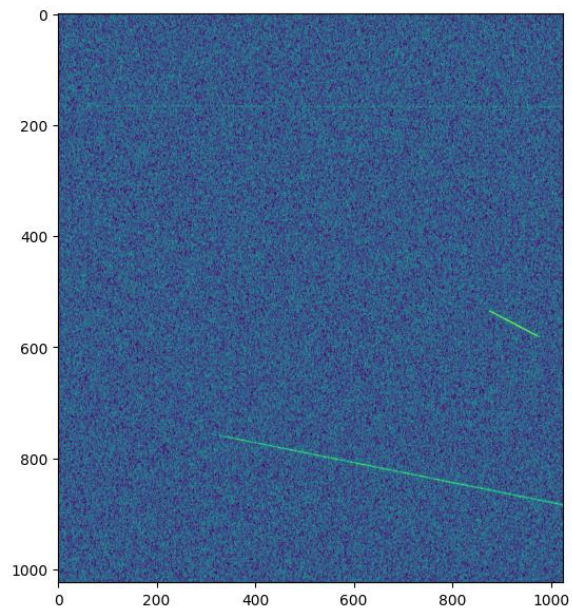
3

2



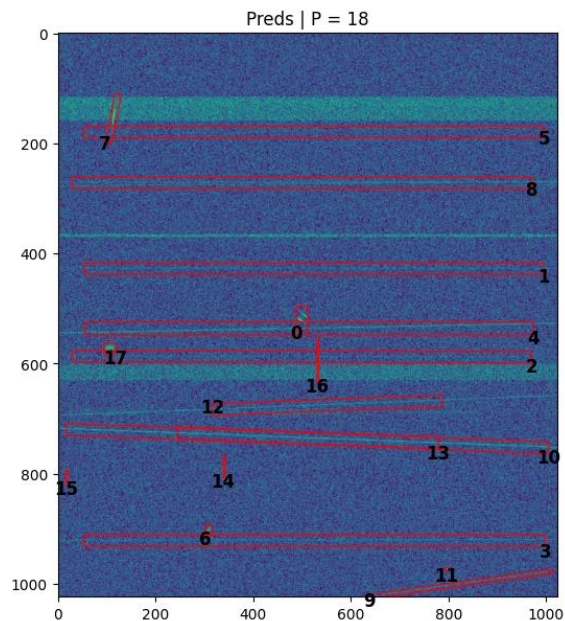
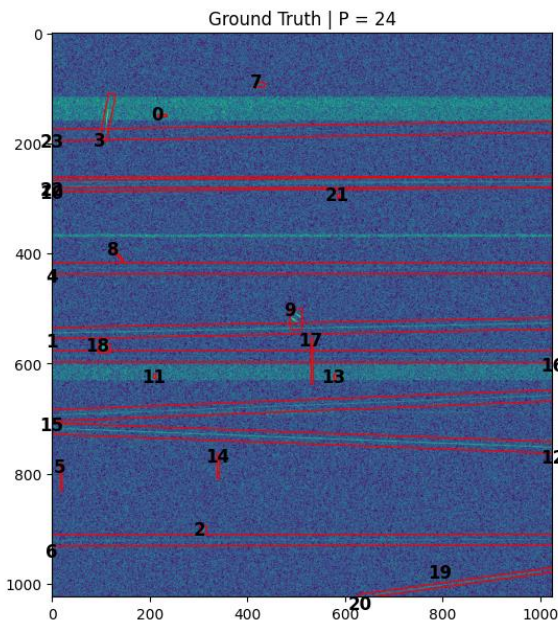
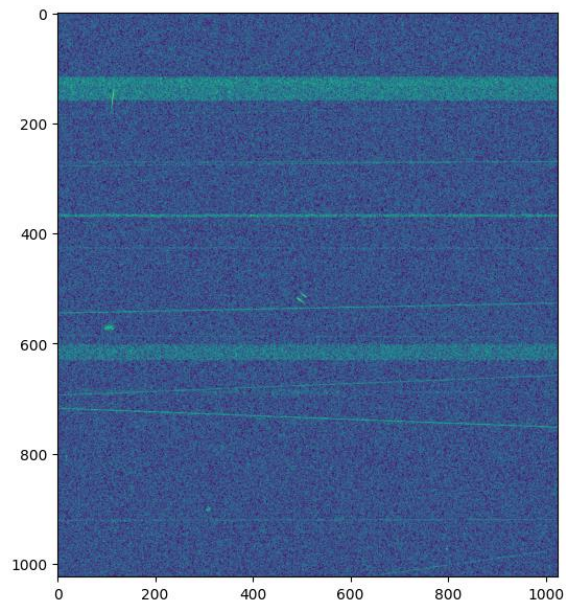
Inférence sur test

seed:8, rec_45715_fs_100.0MSps_nperseg_512.txt



Inférence sur test

seed:9, rec_30824_fs_1000.0MSps_nperseg_1024.txt



AVANTIX

Passage temps contraint

Enumération rapide

Optimisation du modèle

- **Batch processing** : Trouver la bonne valeur du nombre d'image à inférer en parallèle

Optimisation du modèle

- **Quantization** : Passage du FP16 à INT8 (voir FP8)
- **Pruning** : Mise à zéros des poids des neurones peut important. Peut être fait pendant ou après l'apprentissage
- **Knowledge Distillation** : Utiliser les softs output du plus grand modèle afin de compresser le plus gros dans un plus petit

Accélération Hardware

- **TensorRT** : Permet d'accélérer l'inférence sur les GPU NVIDIA en optimisation le graph d'exécution du modèle et/ou de faire de la quantification. Utilisation conjointe avec **ONNX** qui permet de convertir le modèle en un format universel et accepté par TensorRT
- **NVIDIA Triton Inference Server** : Package permettant de faire de l'inférence à hautes performances sur toutes les plateformes

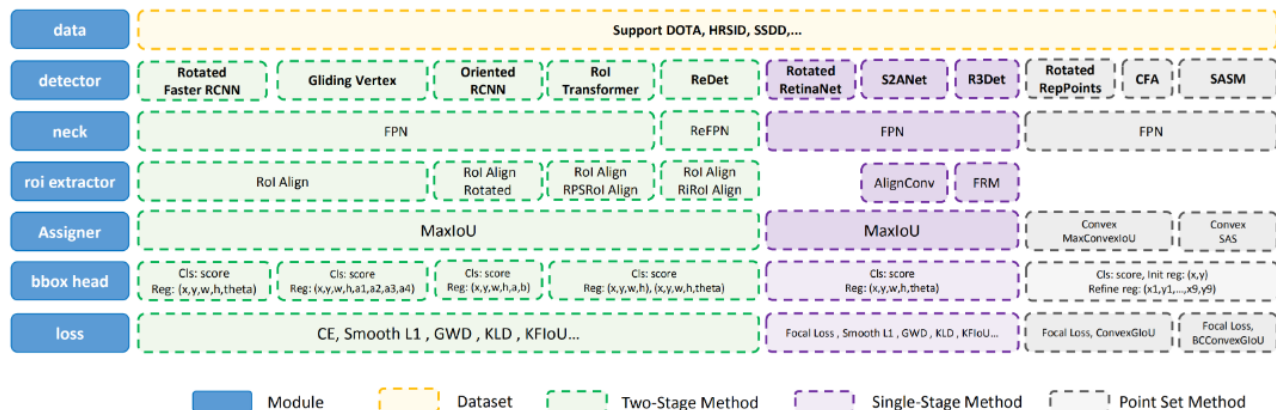
AVANTIX

openMMLab - openMMrotate

MMRotate

[https://github.com/open-mmlab/mmrrotate](https://github.com/open-mmlab/mmrotate)

The module design of MMRotate is as follows:



BENCHMARK AND MODEL ZOO

- [Rotated RetinaNet-OBb/HBB](#) (ICCV'2017)
- [Rotated FasterRCNN-OBb](#) (TPAMI'2017)
- [Rotated RepPoints-OBb](#) (ICCV'2019)
- [Rotated FCOS](#) (ICCV'2019)
- [RoI Transformer](#) (CVPR'2019)
- [Gliding Vertex](#) (TPAMI'2020)
- [Rotated ATSS-OBb](#) (CVPR'2020)
- [CSL](#) (ECCV'2020)
- [R³Det](#) (AAAI'2021)
- [S²A-Net](#) (TGRS'2021)
- [ReDet](#) (CVPR'2021)
- [Beyond Bounding-Box](#) (CVPR'2021)
- [Oriented R-CNN](#) (ICCV'2021)
- [GWD](#) (ICML'2021)
- [KLD](#) (NeurIPS'2021)
- [SASM](#) (AAAI'2022)
- [Oriented RepPoints](#) (CVPR'2022)
- [KFIoU](#) (arXiv)
- [G-Rep](#) (stay tuned)

AVANTIX

TODO

TODO

- BT202c : fs=1GSps nperseg = [128, 1024], 256x1024 (w x h)
- YOLO : Entraîner sur BT300 et mettre en prod le meilleur modèle trouver via HP tuning
- lq2pdw :
 - Découpler inférence et fusion des bbox (inférence par batch)
 - Refact

AVANTIX

MERCI