

AVANTIX

Avancement Simulateur IQ



Binary Segmentation

10/01/2024

© Avantix – for internal use

BT 100 InterMerger

MàJ du scaling

$f_s = 4\text{GSps}$

$N_{iq} = 131072 * 128$

$N_{PERSEGS} = [128, 256, 512, 1024, 2048, 8192, 32768, 131072]$

Tous les spectrogrammes sans overlap et sans fenêtrage :

- $x = x.reshape(nperseg, -1)$
- $X = fft(x)$
- $sxx = \frac{|X|^2}{nperseg * \sigma_{noise}^2}$

Seuil de binarisation « optimal » = 0 dB

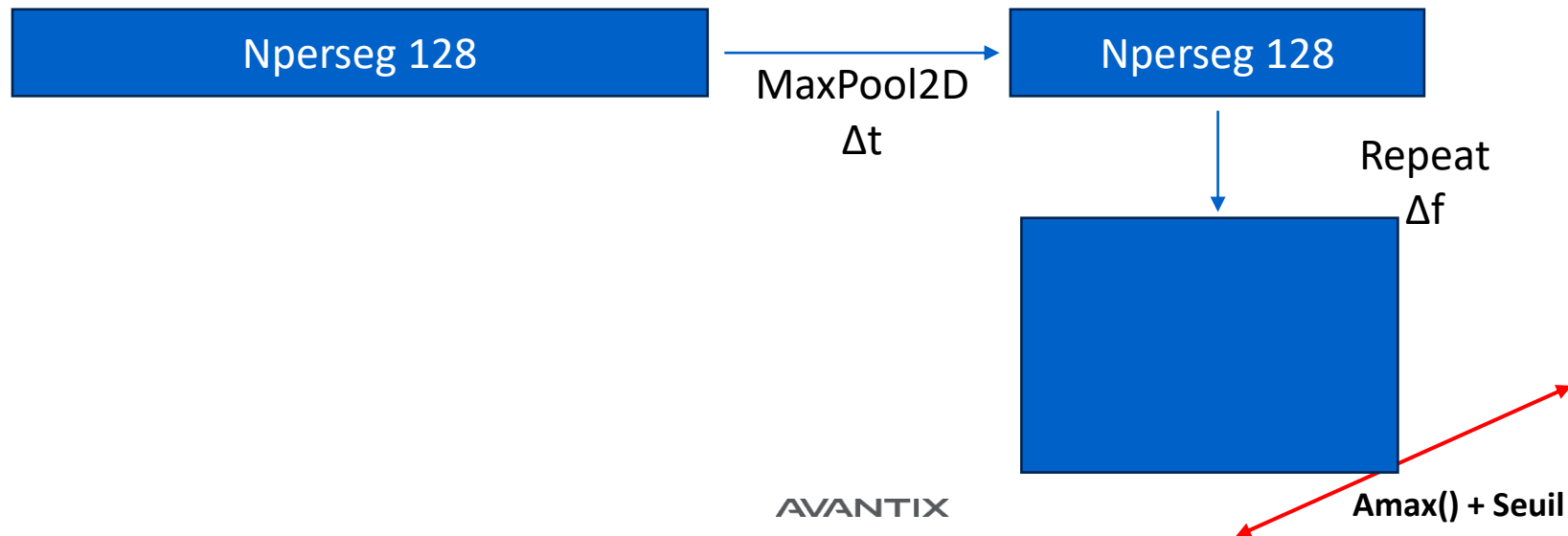
Résolution du masque : ~1MHz en freq et ~1μs en temps

Définition du masque : 4096, 4096

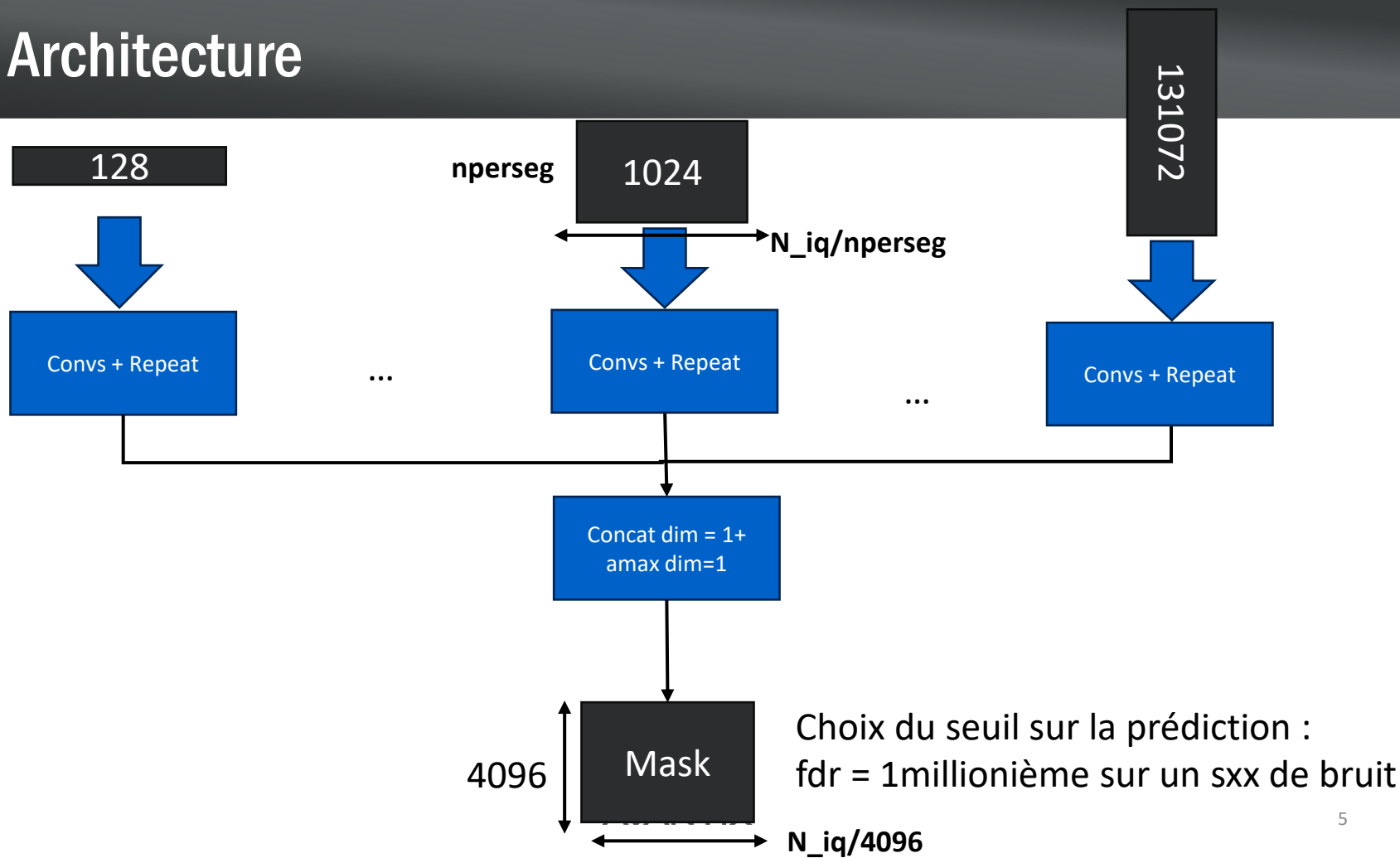
BT 100 InterMerger

Construction du masque :

1. Sur noiseless_sxx à tous les nperseg :
 - MaxPool sur la dimension la plus grande jusqu'à atteindre 4096
 - UpSampling en répétant sur l'autre dimension jusqu'à atteindre 4096
2. $\text{AMAX}(\text{noiseless_sxx_resized}) > -3 \text{ dB}$



Architecture



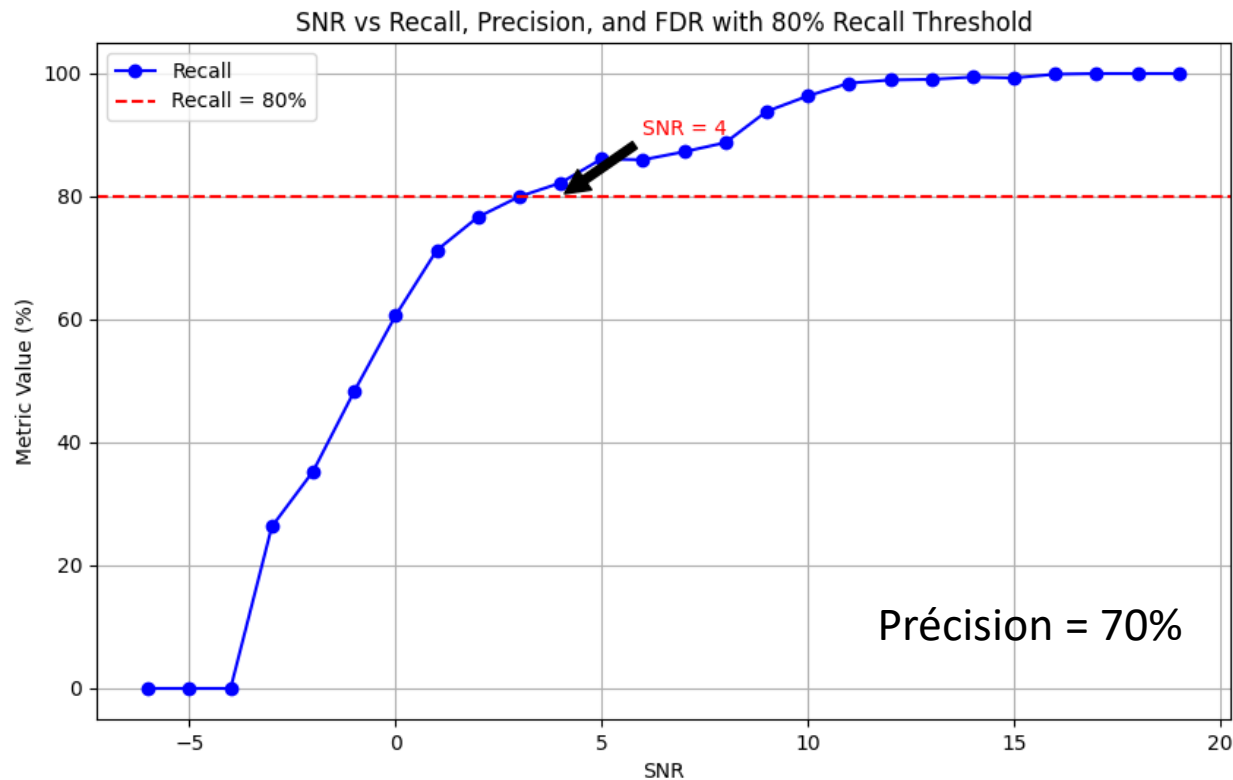
Torch vs TensorRT

Torch (nn Module) : 3,9

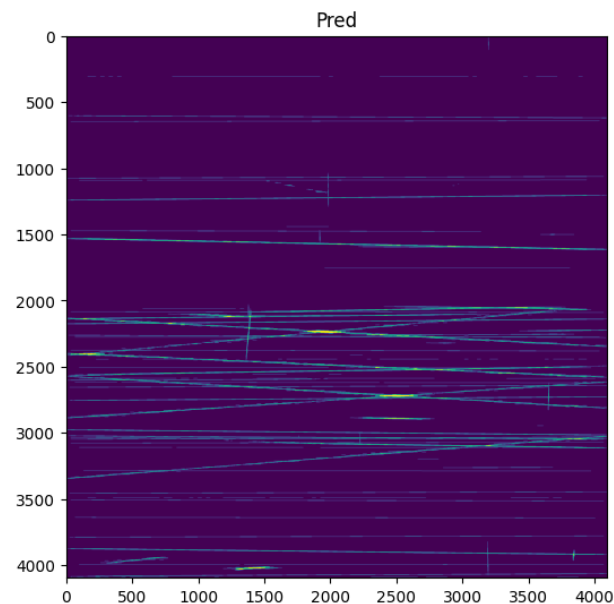
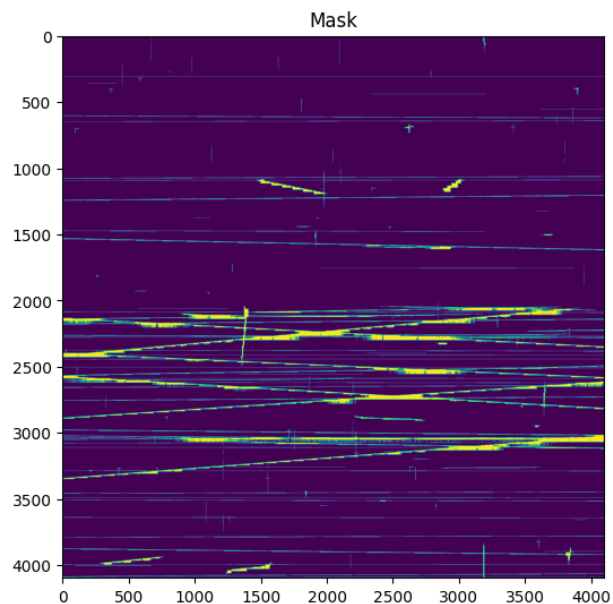
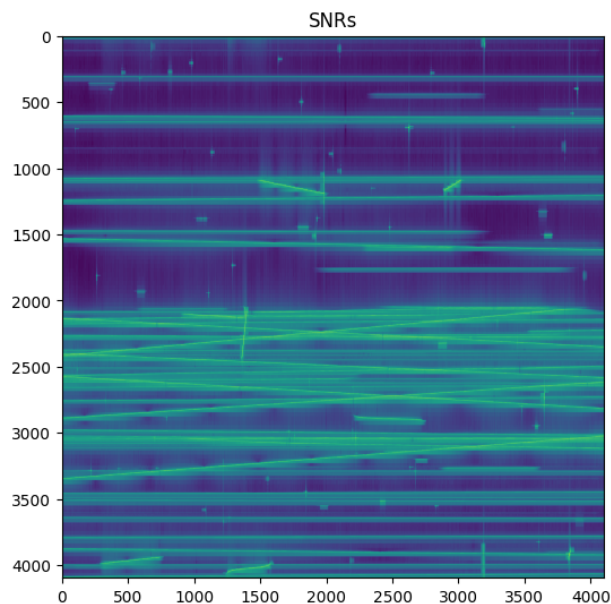
Avec TensorRT : 4,5

Hypothèse : TensorRT a du mal à optimiser les modèles à plusieurs entrées et/ou ayant des couches de repeat

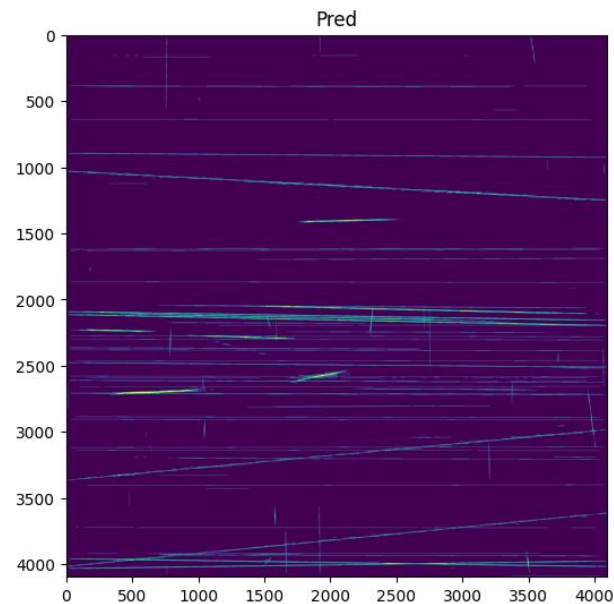
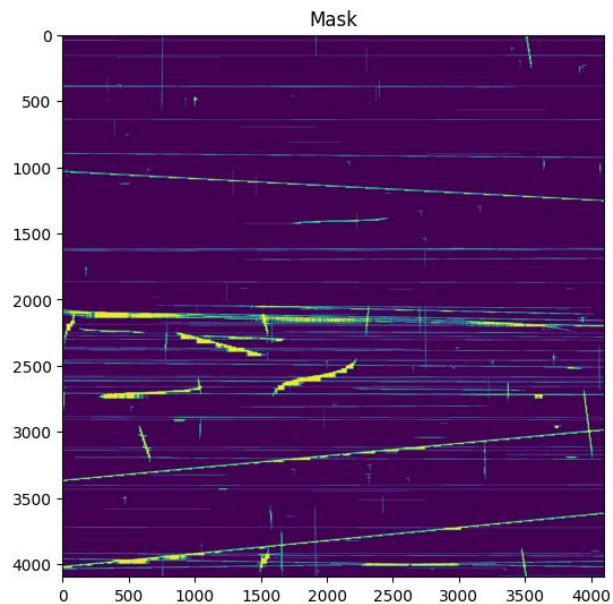
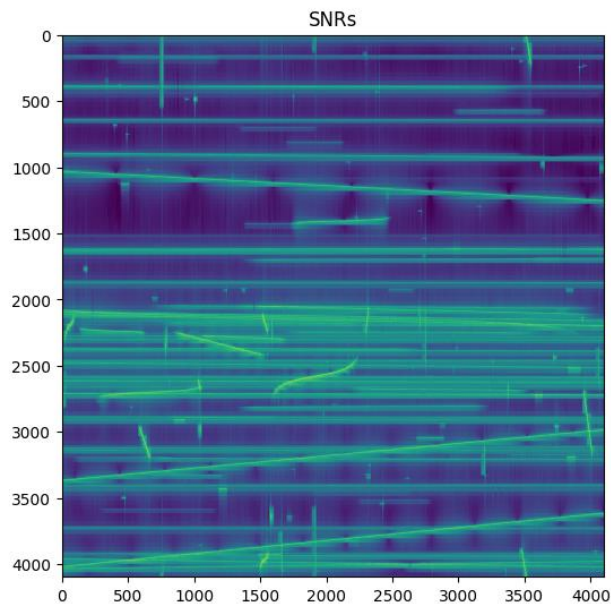
Num_kernels = 1 Résultats inférence



Num_kernels = 1 Résultats inférence

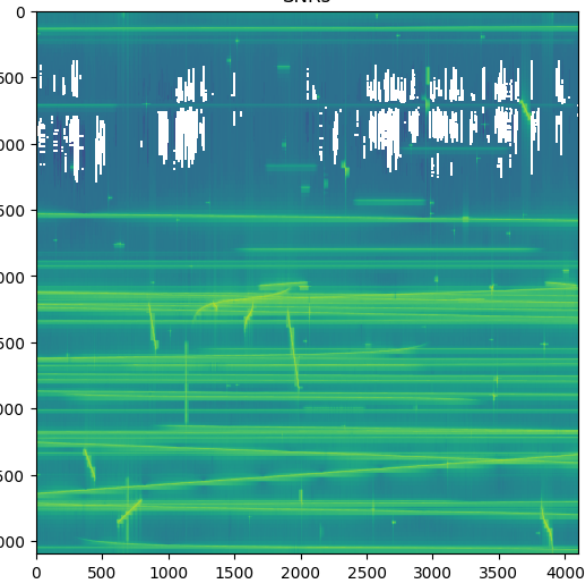


Num_kernels = 1 Résultats inférence

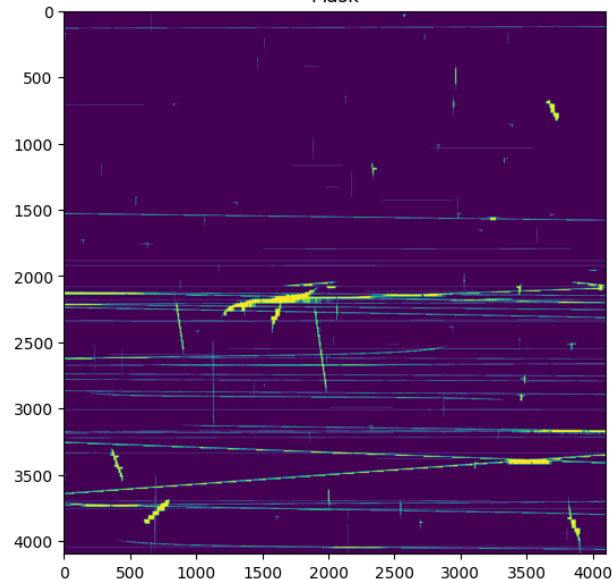


Num_kernels = 1 Résultats inférence

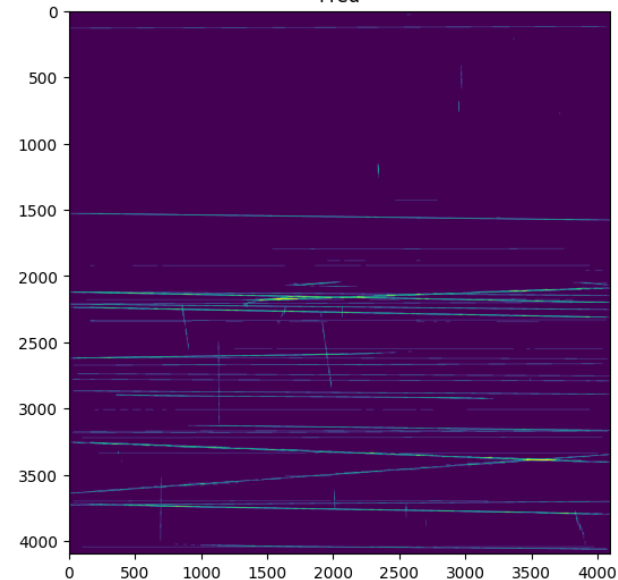
SNRs



Mask

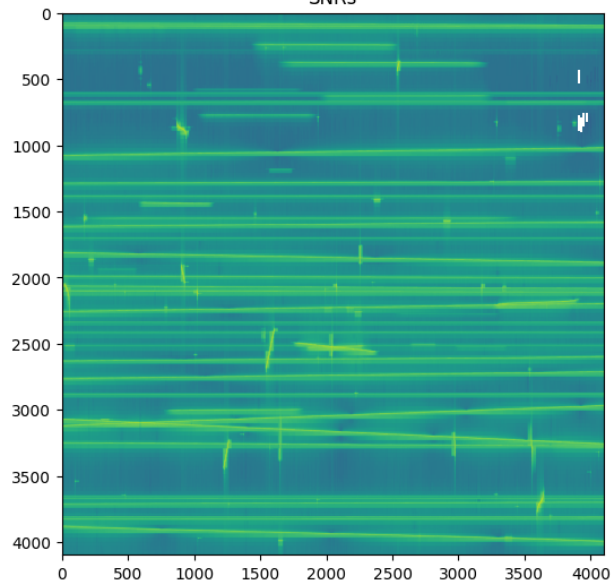


Pred

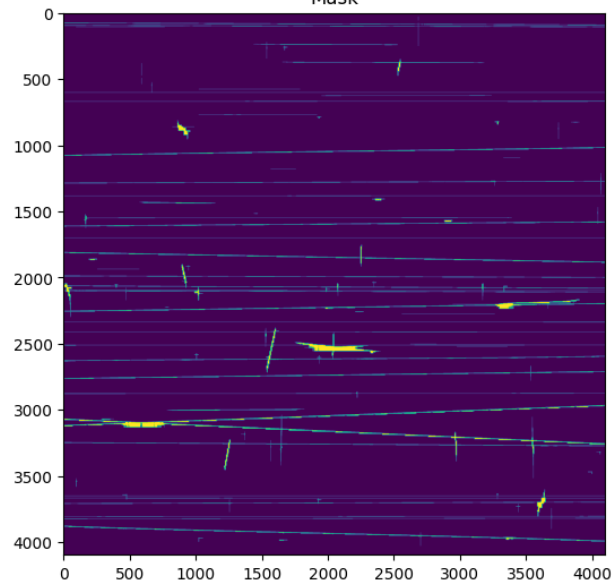


Num_kernels = 1 Résultats inférence

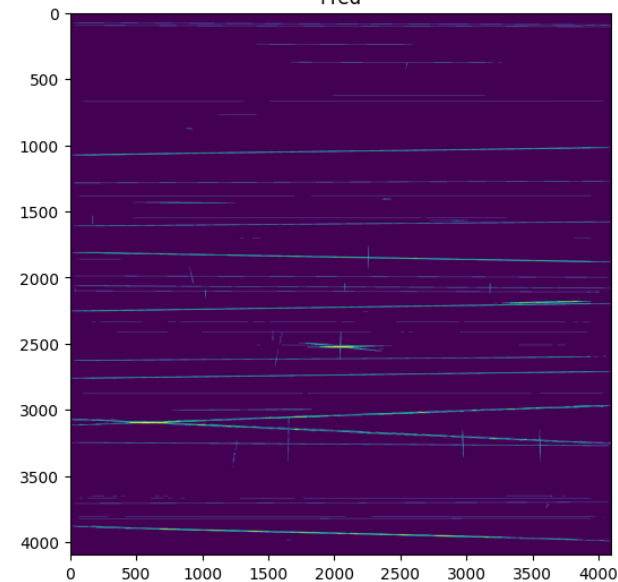
SNRs



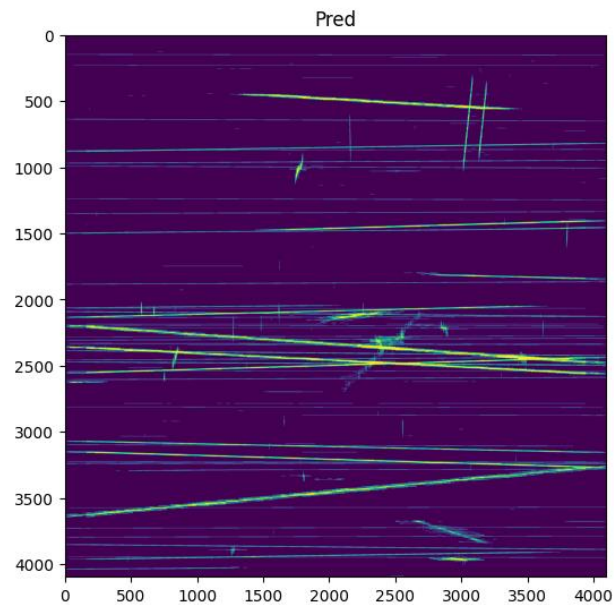
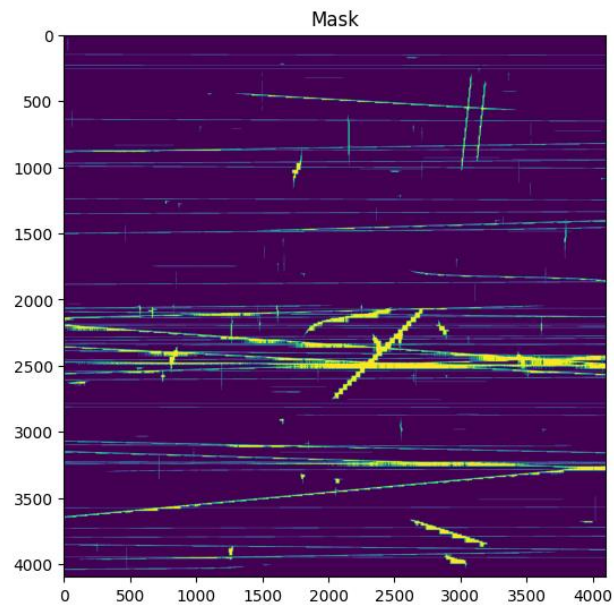
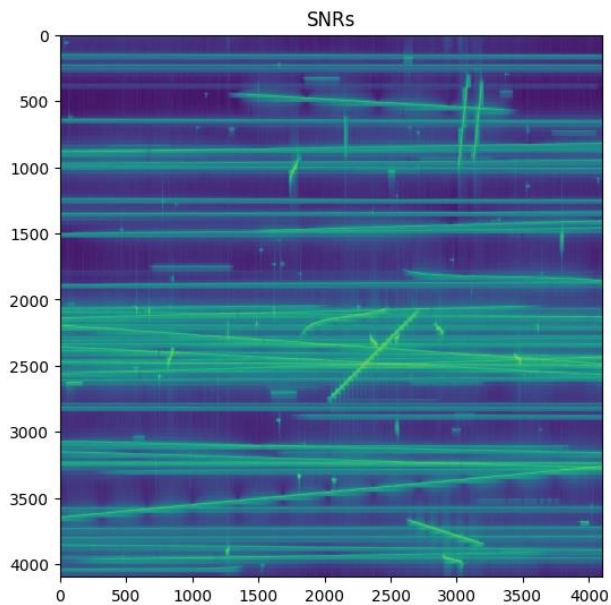
Mask



Pred



Num_kernels = 1 Résultats inférence





Experimentation input_shape

Ratio temps réel modèle InterMerger

L'objectif était de savoir si le RTR dépend du type de formatage :
concaténation sur batch vs temporelle

Résultat : aucune différence

Input shape avec num_kernels = 8	RTR
(16, 1, 128, 131072) Concaténation batch	1,5
(1, 1, 128, 16*131072) Concaténation tempo	1,5

AVANTIX

MERCI