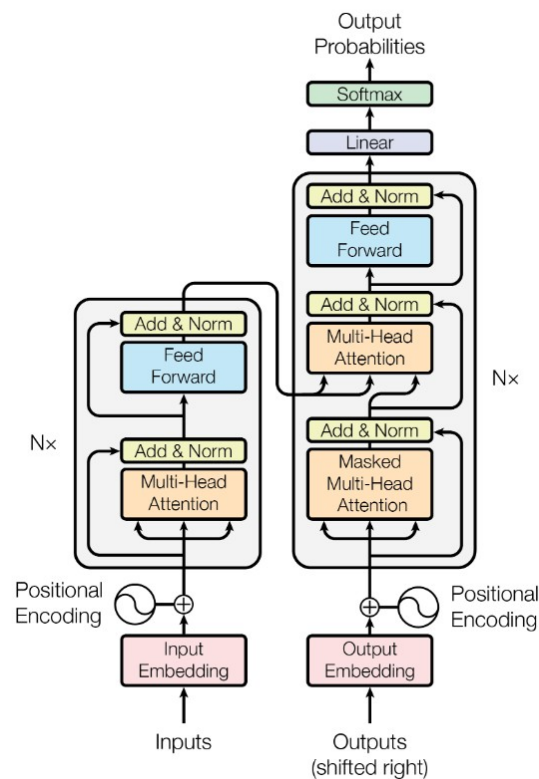


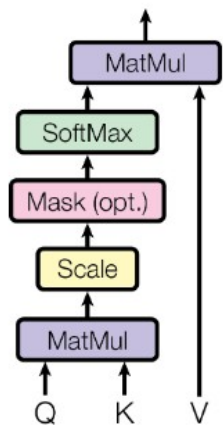


Visual Transformer

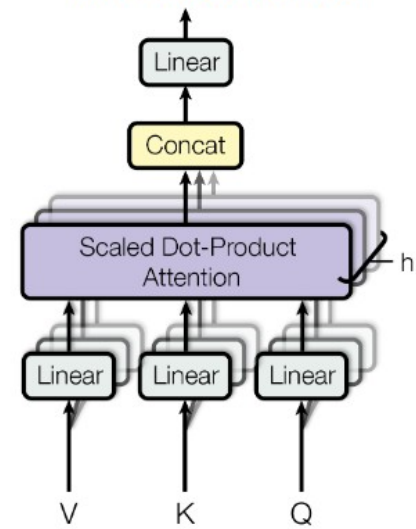
Attention is all you need (Vaswani 2017)



Scaled Dot-Product Attention

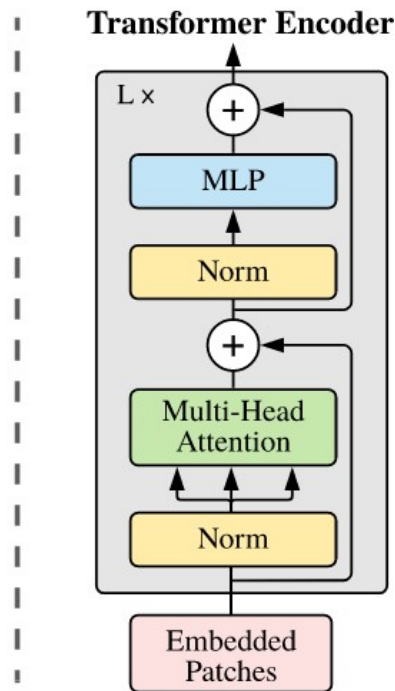
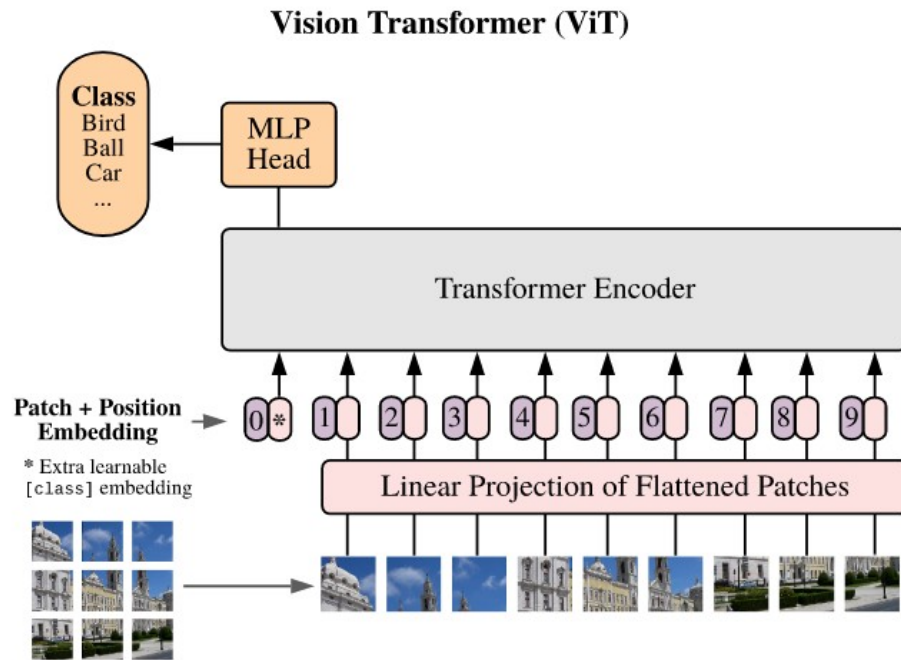


Multi-Head Attention

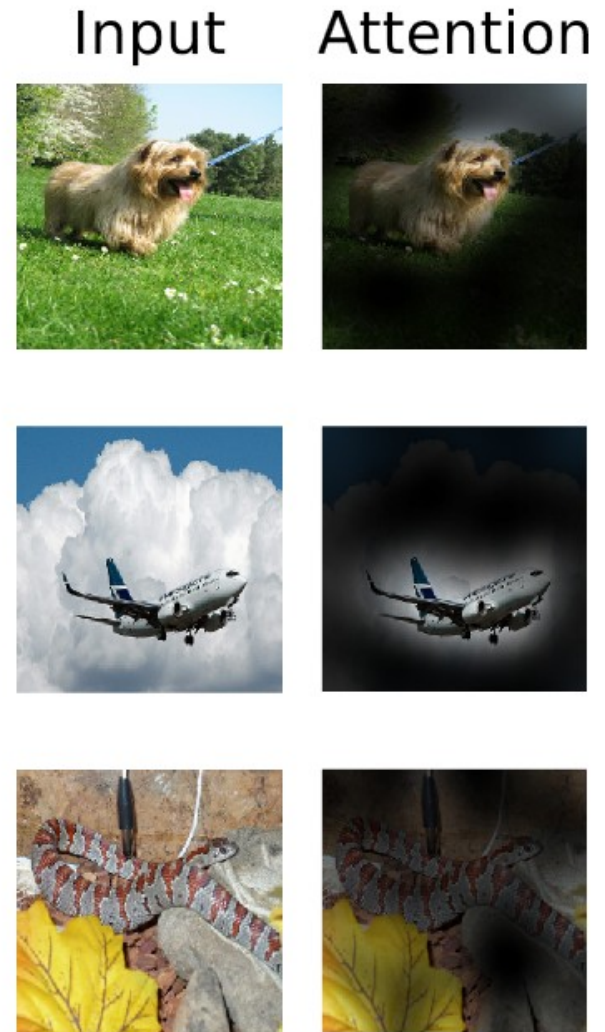
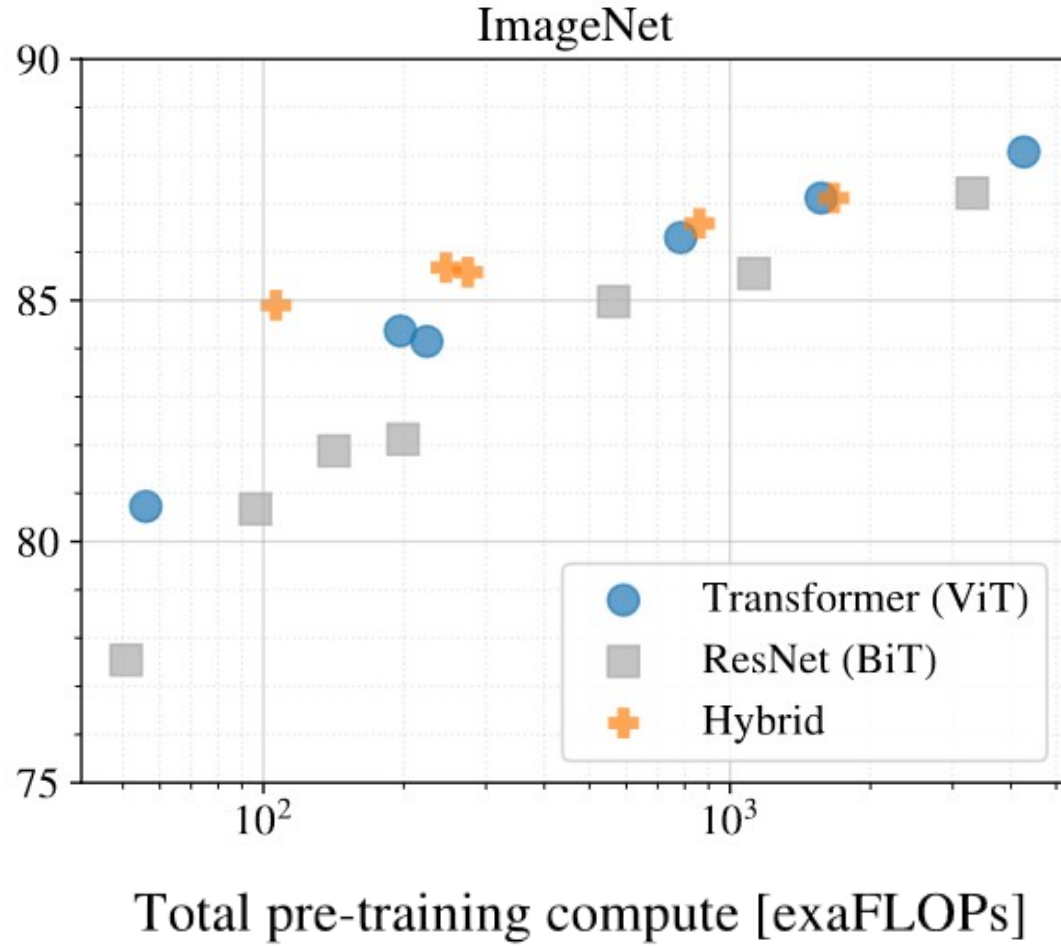




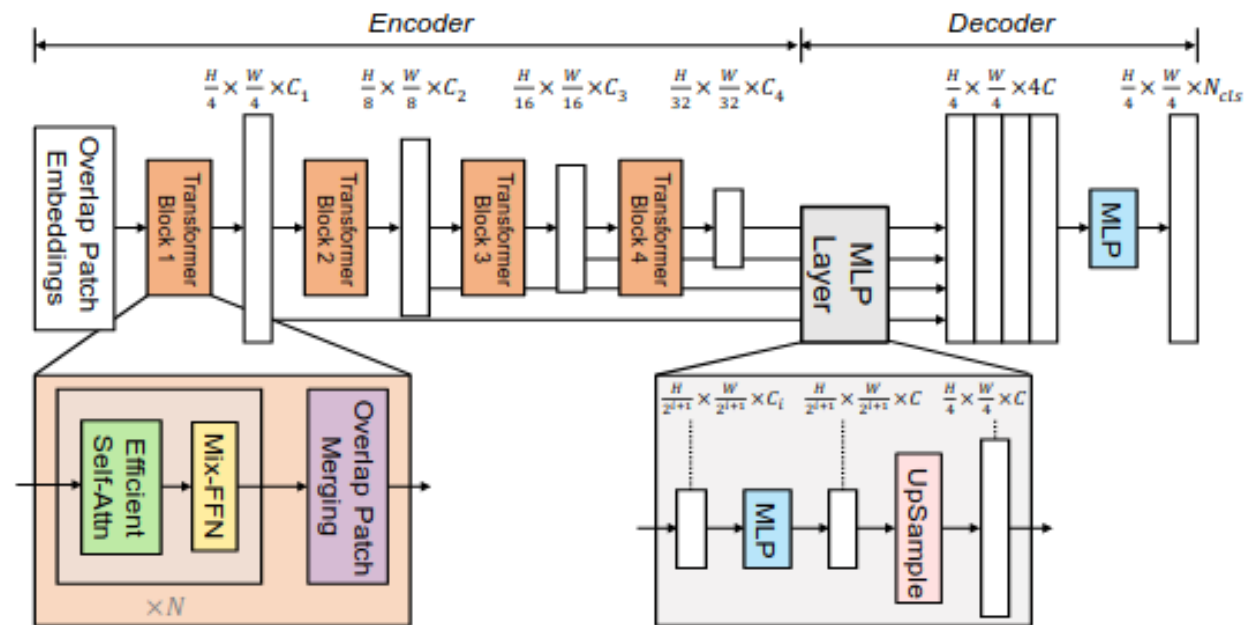
An Image is Worth 16x16 Words: Transformers for Image Recognition at Scale (Dosovitskiy 2020)



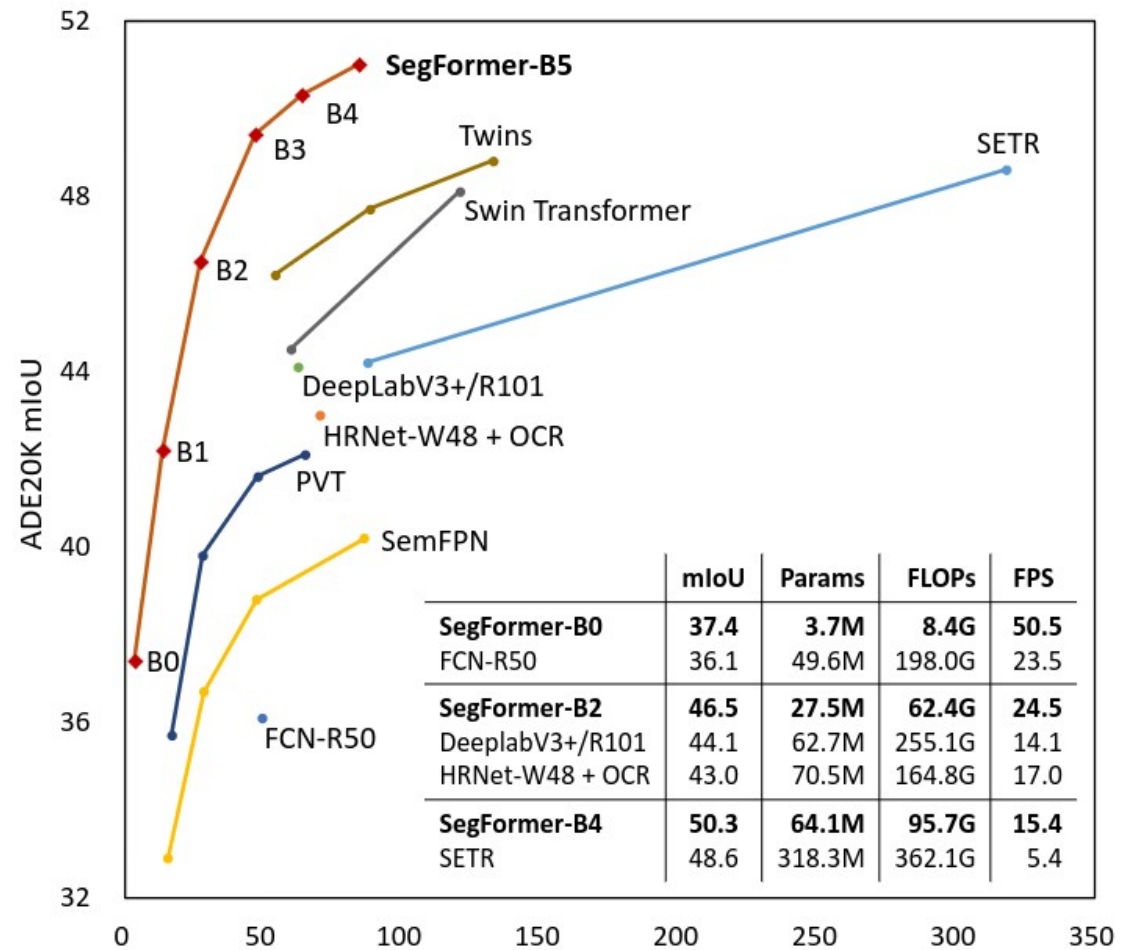
- Image comme ensemble de patch (16x16)
- Ajout "faux" token pour classif (BERT, Devlin2019)
- 1D conv sur patch + position : ordre rasterization

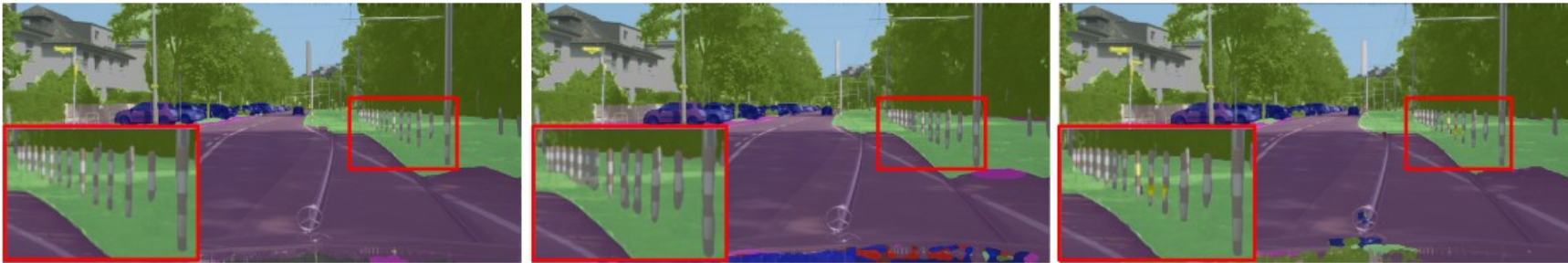


SegFormer: Simple and Efficient Design for Semantic Segmentation with Transformers (Xie 2021)



- Retire l'encodage de la position
- Utilise une fenêtre glissante pour extraire les patches





SegFormer

SETR

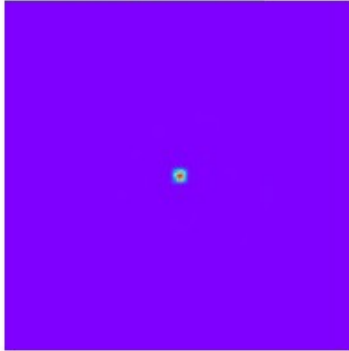
DeepLabV3+

Pourquoi ?

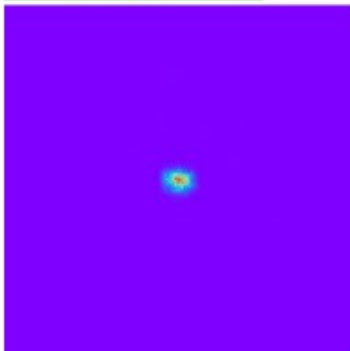
DeepLabv3+



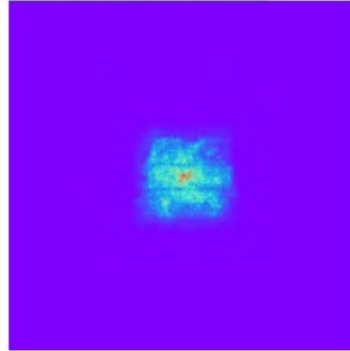
Stage-1



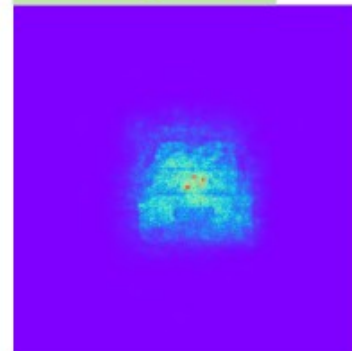
Stage-2



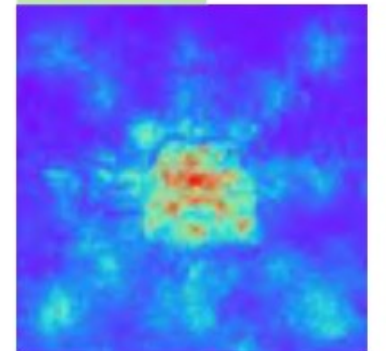
Stage-3



Stage-4



Head

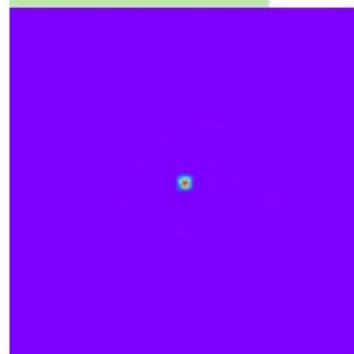


Pourquoi ?

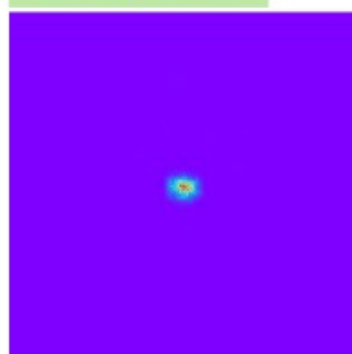
DeepLabv3+



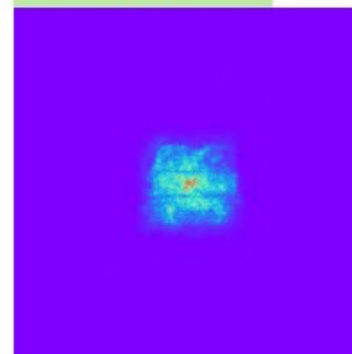
Stage-1



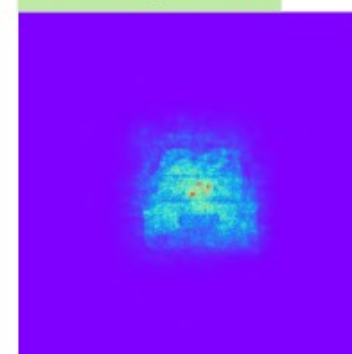
Stage-2



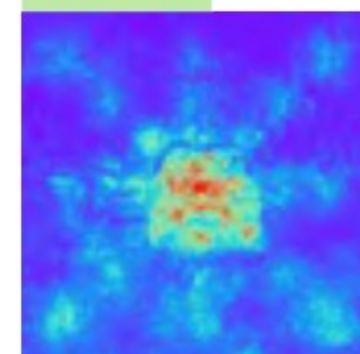
Stage-3



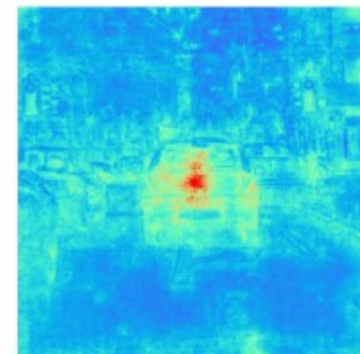
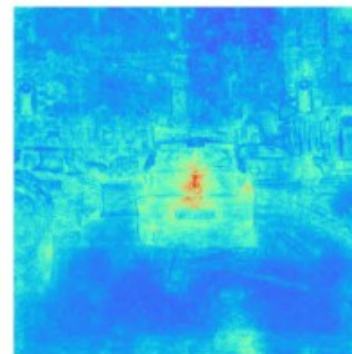
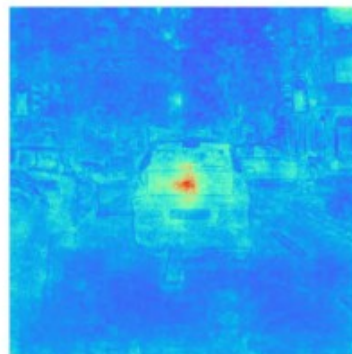
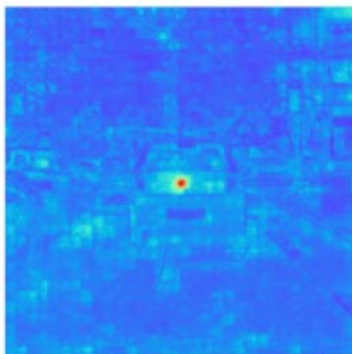
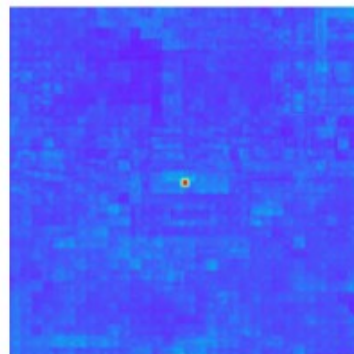
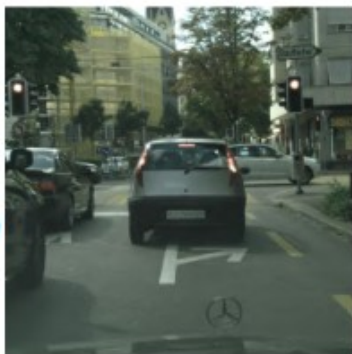
Stage-4



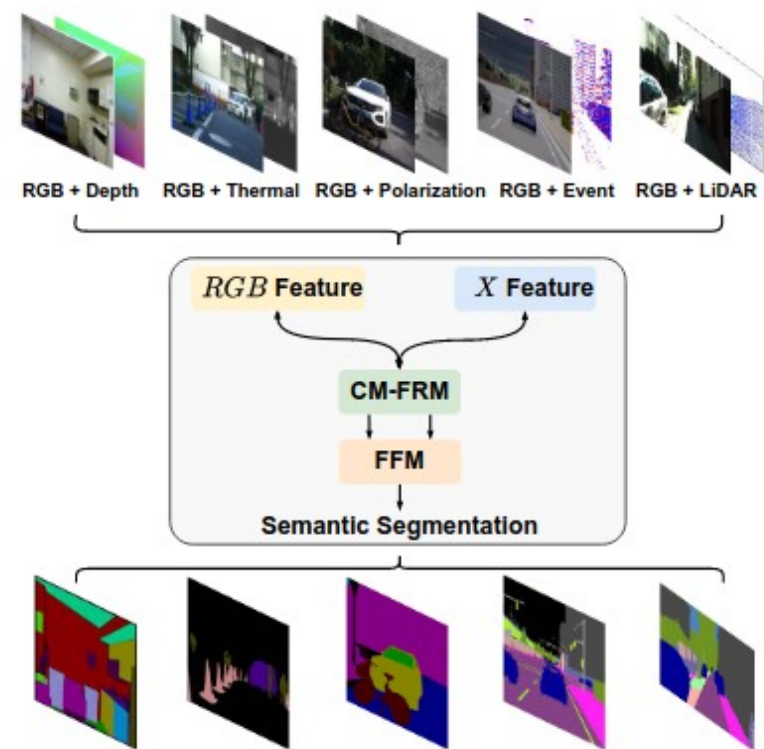
Head



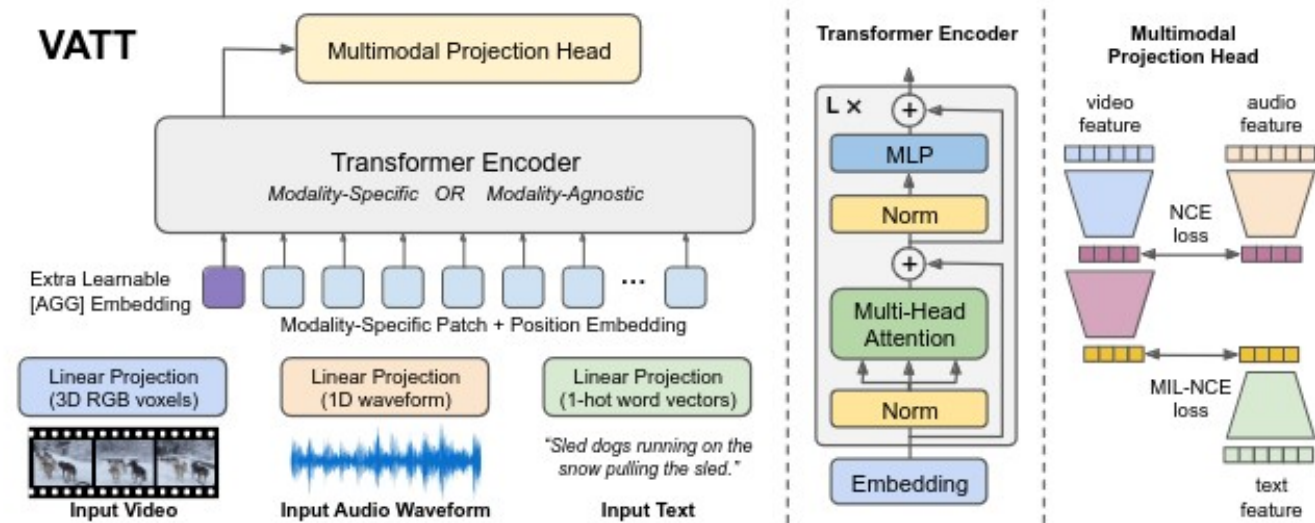
SegFormer



Approche multimodal



CMX (Zhang 22)



VATT (Akbari 21)

Ref

"Attention Is All You Need", Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N. Gomez, Lukasz Kaiser, Illia Polosukhin - NIPS2017

"An Image is Worth 16x16 Words: Transformers for Image Recognition at Scale", Alexey Dosovitskiy, Lucas Beyer, Alexander Kolesnikov, Dirk Weissenborn, Xiaohua Zhai, Thomas Unterthiner, Mostafa Dehghani, Matthias Minderer, Georg Heigold, Sylvain Gelly, Jakob Uszkoreit, Neil Houlsby - ICLR2021

"SegFormer: Simple and Efficient Design for Semantic Segmentation with Transformers", Enze Xie, Wenhai Wang, Zhiding Yu, Anima Anandkumar, Jose M. Alvarez, Ping Luo - NeurIPS 2021.

"CMX: Cross-Modal Fusion for RGB-X Semantic Segmentation with Transformers", Jiaming Zhang, Huayao Liu, Kailun Yang, Xinxin Hu, Ruiping Liu, Rainer Stiefelhagen

"VATT: Transformers for Multimodal Self-Supervised Learning from Raw Video, Audio and Text", Hassan Akbari, Liangzhe Yuan, Rui Qian, Wei-Hong Chuang, Shih-Fu Chang, Yin Cui, Boqing Gong – NeurIPS 2021