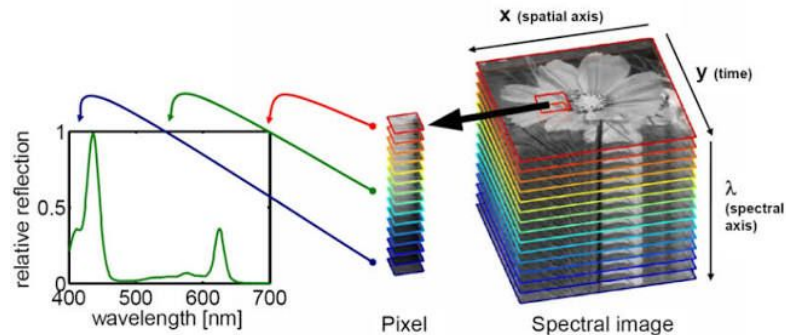


Hyperspectral Foundation-Model Feature Upsampling for Unsupervised Unmixing: Training NAF on Frozen DOFA Features

TABOURI Lukas, ESSAADANI Obeye, EL GOUCH Hussein, BALI Samy

Imagerie hyperspectrale & démixage spectral

- Un pixel hyperspectral = un spectre quasi-continu sur des centaines de bandes étroites.
- Compromis : forte dimension spectrale, mais résolution spatiale grossière.
- Conséquence : pixels mixtes mélange de plusieurs matériaux purs (endmembers).



Problème inverse **modèle de mélange linéaire**

$$Y = AS + N, \quad A \in \mathbb{R}^{B \times R}, \quad S \in \mathbb{R}^{R \times P}$$

$$S \geq 0 \quad (\text{ANC}) \quad \mathbf{1}^\top S = \mathbf{1}^\top \quad (\text{ASC})$$

Colonnes de A = spectres ; S = abondances par pixel. Exactement une NMF.

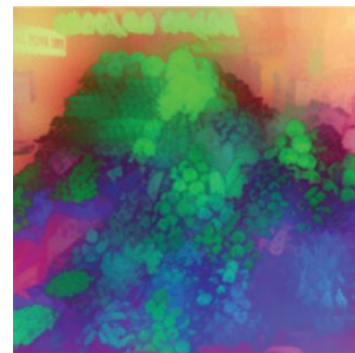
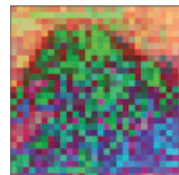
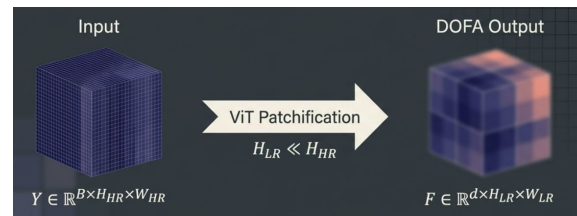
Foundation models & le fossé de résolution

DOFA : hypernetwork conditionné par la longueur d'onde qui génère les poids du patch-embedding → un seul ViT pour un nombre arbitraire de bandes.

Comme tout ViT, DOFA patchifie l'image → carte de features à résolution réduite :

$$F \in \mathbb{R}^{d \times H_{LR} \times W_{LR}}, \quad H_{LR} \ll H_{HR}$$

Sémantiquement riches mais **spatialement grossières** → inutilisables telles quelles pour une tâche par-pixel comme le démelange.



NAF

VFM : La résolution spatiale est réduite, on a une perte de précision. Comment y remédier ?

NAF : premier module d'upsampling totalement agnostique (indépendant du VFM utilisé) capable de surpasser les méthodes spécifiques. Il fonctionne en "**zero-shot**", c'est-à-dire qu'une fois entraîné, il peut s'appliquer instantanément à n'importe quel modèle de vision sans aucun réentraînement

Pourquoi NAF ? Filtres classiques : résultats flous et indépendants du contexte

Upsamplers modernes : JAFAR, FeatUp sont spécifiques à chaque VFM, il faut donc réentraîner le module d'upsampling pour chaque VFM

OBJECTIF & CONTRIBUTION

Adapter NAF à l'hyperspectral et l'entraîner sur features DOFA gelées

Ce qui est gelé

DOFA + l'adaptateur spectral
SpectralEarth restent figés.

Ce qui est entraîné

Seul le module NAF (encodeur
de guidage + attention).

Objectif sans label

Self-distillation : un seul terme
MSE en espace de features.

Pipeline entièrement label-free → directement compatible avec un démixage non-supervisé.

Trois composants, deux gelés, un entraîné



DOFA extracteur de features (gelé)

Mappe l'HSI vers la carte basse résolution F, qui fournit les valeurs V_{LR} de l'attention.



Adaptateur spectral SpectralEarth (gelé)

Mappe B bandes (variables selon capteur) vers un tenseur fixe de 128 canaux → voie de guidage agnostique au capteur.



NAF upsampler image-guidé (entraîné)

Agnostique au foundation model : upsampling vu comme un filtrage adaptatif appris, guidé par l'image HR.

Attention de voisinage cross-échelle + RoPE

Pour chaque position HR p , on attend sur un petit voisinage $\mathcal{N}(p)$ de sa projection basse résolution :

$$F_{up}(p) = \sum_{q \in \mathcal{N}(p)} \alpha_{pq} V_{LR}(q)$$

$$\alpha_{pq} = \text{softmax}_{q \in \mathcal{N}(p)} \left(\frac{Q_{HR}(p)^\top K_{LR}(q)}{\sqrt{C}} \right)$$

$\mathbf{Q}, \mathbf{K}$ portent RoPE $\rightarrow \alpha$ dépend du déplacement relatif $p-q$.

Valeurs = features DOFA gelées : l'upsampler n'apprend que « où regarder », jamais à modéliser la distribution des features.

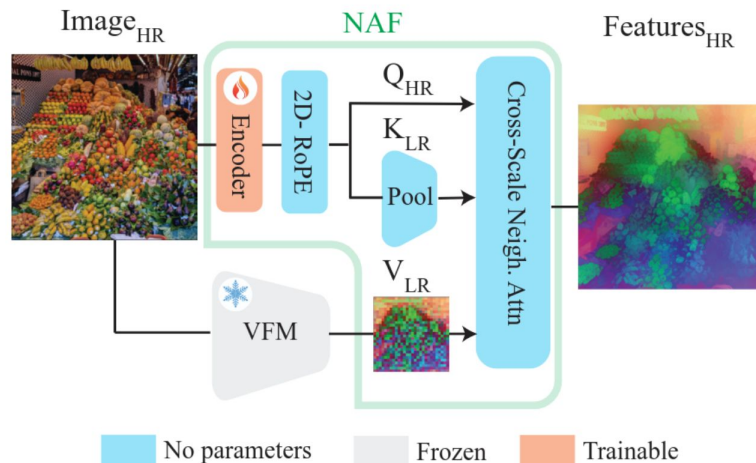
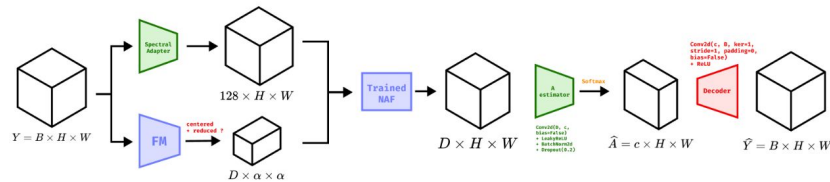


Figure 2: Naf components



- 1 Extraction gelée.** DOFA $\rightarrow$ carte de features = valeurs V_{LR} de l'attention.
- 2 Encodage de guidage.** Adaptateur $\rightarrow$ 128 canaux ; encodeur dual-branch $\rightarrow$ carte de guidage HR.
- 3 Construction Q / K.** RoPE sur le guidage $\rightarrow Q_{HR}$; average-pooling LR $\rightarrow K_{LR}$.
- 4 Attention cross-échelle.** Pour chaque position HR, attention sur un voisinage LR $\rightarrow F_{up}$.

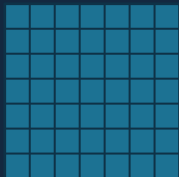
Seuls les paramètres θ de NAF (encodeur de guidage + attention) sont mis à jour ; DOFA et l'adaptateur spectral restent gelés.

Self-distillation des features : protocole 7 $\rightarrow$ 14

Un seul terme MSE en espace de features

$$\mathcal{L}(\theta) = \| \text{NAF}_{\theta}(\downarrow F_{\text{HR}}; Y) - F_{\text{HR}} \|_2^2$$

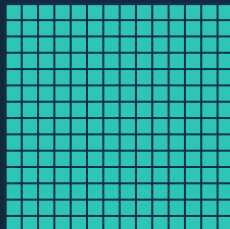
DOFA (14×14) = professeur ; cible figée, aucun gradient vers DOFA.



F_LR 7×7



NAF ($s = 2$)



F_up 14×14

Pourquoi sous-échantillonner ?

Aucune feature HR n'existe au-delà de DOFA $\rightarrow$ on simule l'upsampling en partant d'une copie sous-échantillonnée.

Opérateur image-guidé & résolution-agnostique $\rightarrow$ à l'inférence on dépasse 14×14.

Zéro label $\rightarrow$ démélange aval non-supervisé.

Données, backbones gelés & métriques

Données

- Entraînement : EnMAP (SpectralEarth).
- Évaluation : images avec endmembers & abondances de référence.

Backbones gelés

- DOFA-Large (ViT-L) $\rightarrow d = 1024$.
- Adaptateur spec_ViTb_mae $\rightarrow 128$ canaux.
- Seuls les poids de NAF sont entraînables.

Métriques d'évaluation

MSE de reconstruction entre F_{up} et la cible DOFA (7 $\rightarrow$ 14) **SAD** $\downarrow$ angle spectral endmembers estimés / référence.

Abundance RMSE/NMSE $\downarrow$ cartes d'abondance ; **PCA de F vs F_{up}** diagnostic qualitatif (détail récupéré, pas lissé).

Baseline : upsampling bilinéaire.

NAF vs bilinéaire — mêmes features, même tête

-27%

SAD endmembers

0.099 → 0.072

≈ 3%

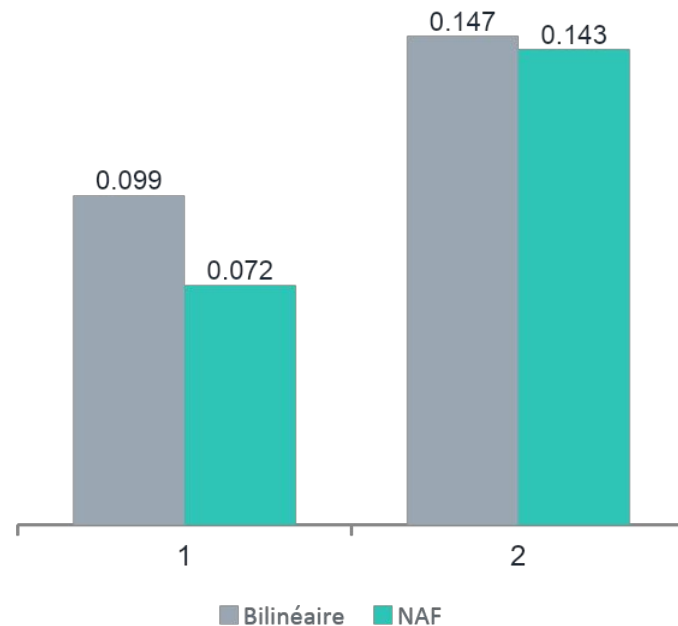
NMSE abundances

0.147 → 0.143

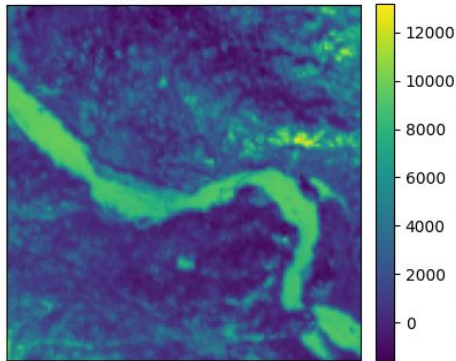
Gain net et cohérent sur les spectres ; quasi-égalité sur les abundances, dominée par un seul endmember (EM1).

Détail par endmember

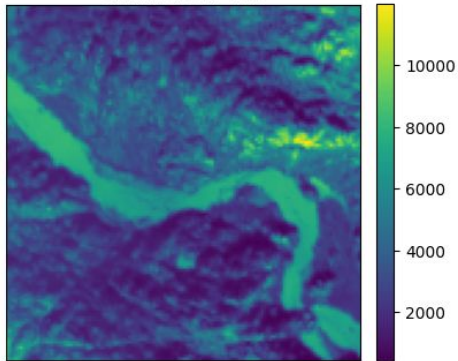
	EM1	EM2	EM3	EM4
SAD · Bilin	0.14	0.14	0.12	0.05
SAD · NAF	0.10	0.14	0.10	0.04
NMSE · Bilin	0.99	0.05	0.16	0.12
NMSE · NAF	0.98	0.06	0.15	0.10



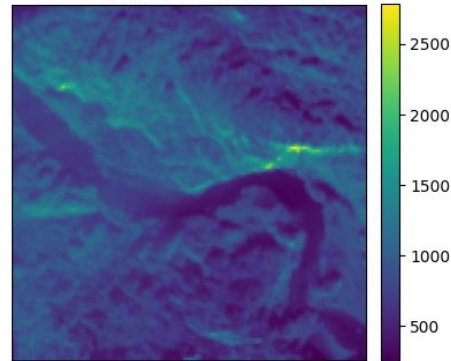
Channel 0 / 202



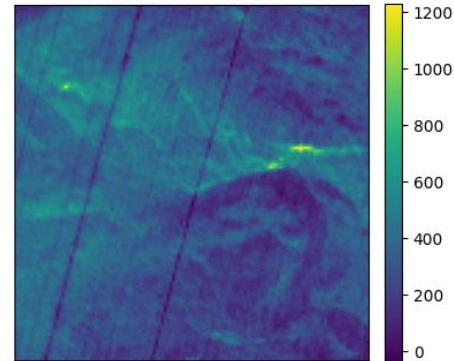
Channel 67 / 202



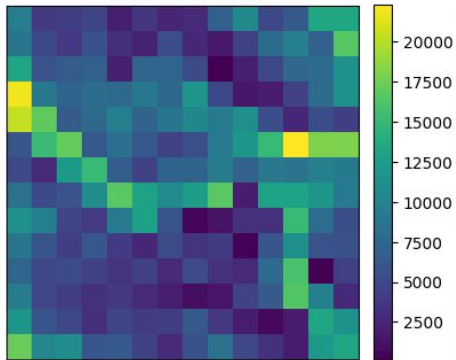
Channel 134 / 202



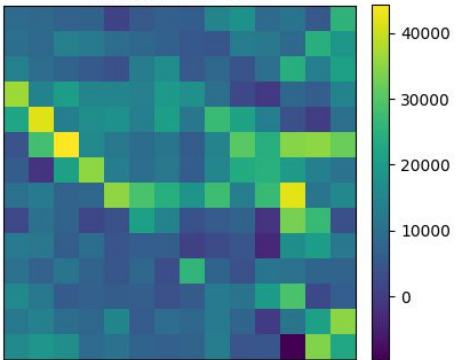
Channel 201 / 202



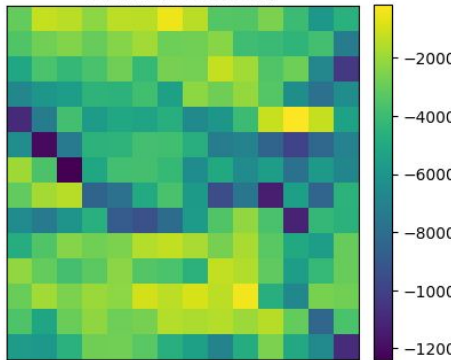
Channel 0 / 1024



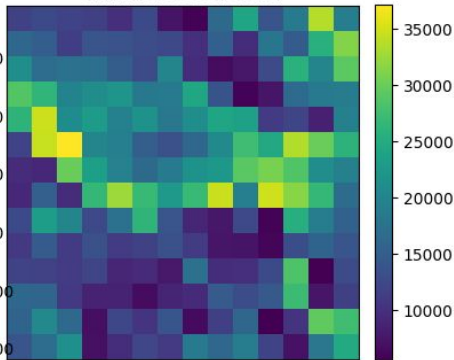
Channel 341 / 1024



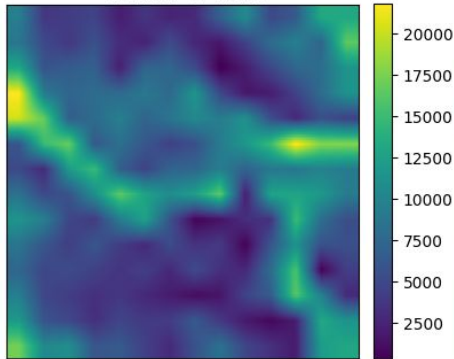
Channel 682 / 1024



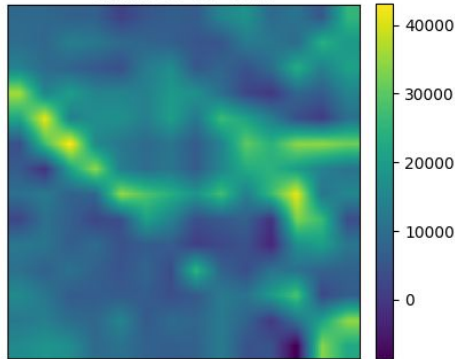
Channel 1023 / 1024



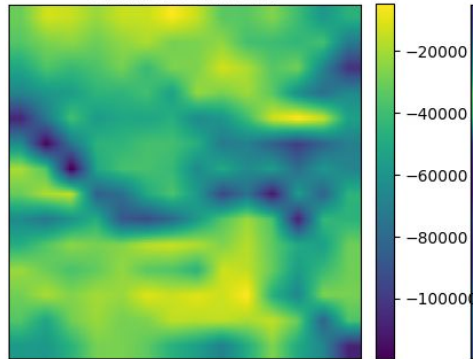
Channel 0 / 1024



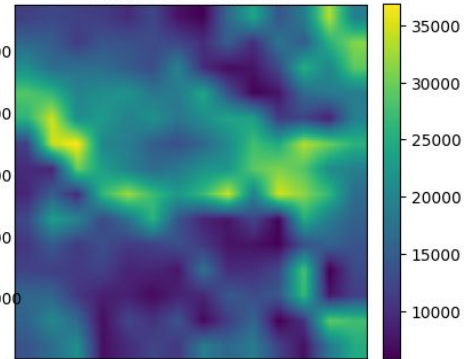
Channel 341 / 1024



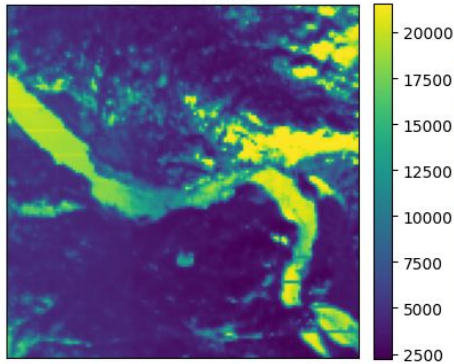
Channel 682 / 1024



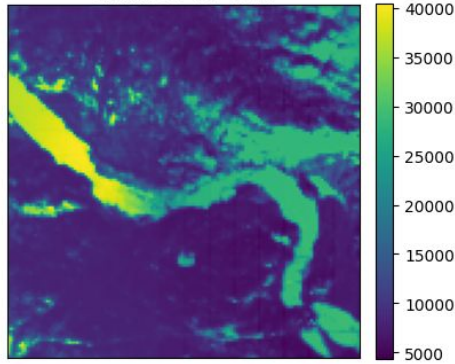
Channel 1023 / 1024



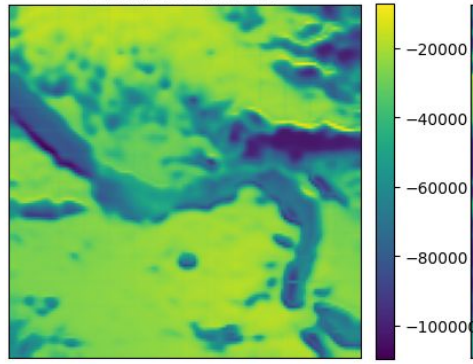
Channel 0 / 1024



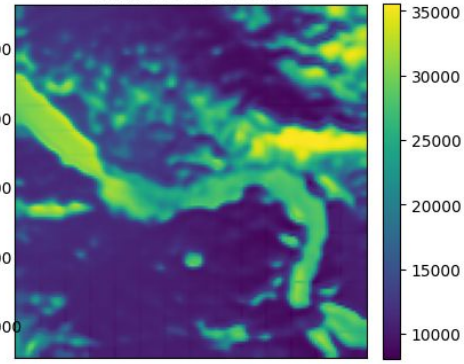
Channel 341 / 1024



Channel 682 / 1024



Channel 1023 / 1024

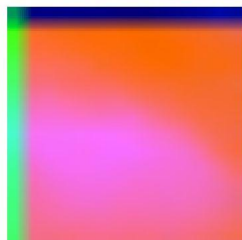
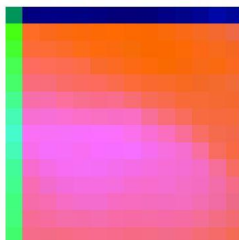
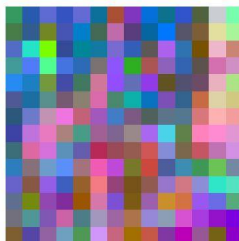
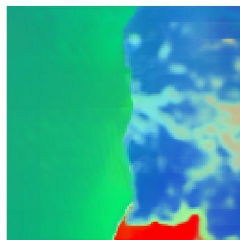
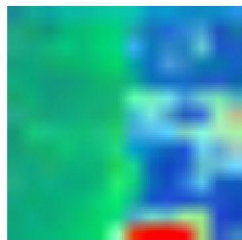
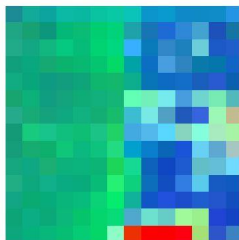
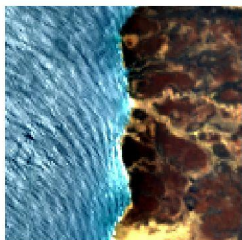
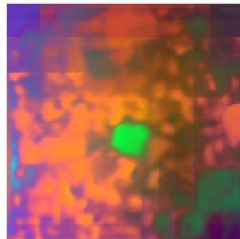
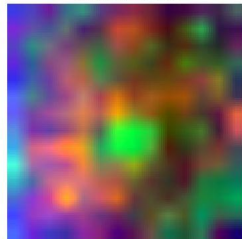
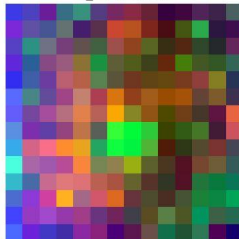


EnMAP RGB

F_LR (14x14)

Bilinéaire (128)

NAF (128)



Cartes plus nettes & spectres resserrés

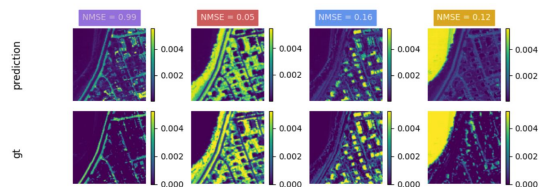


Figure 5: abundance DOFA+ bilinear

Abundance estimation, NMSE = 0.143

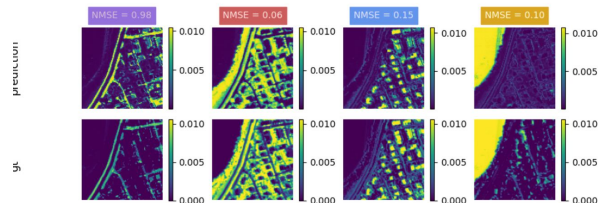


Figure 6: abundance NAF

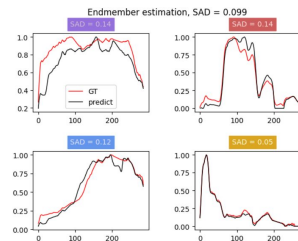


Figure 3: endmember DOFA+ bilinear

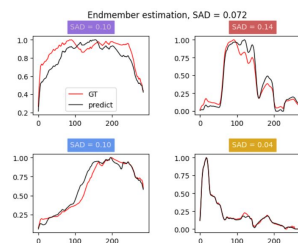


Figure 4: endmember NAF

Cartes d'abondance

Réseau routier & blocs de bâtiments nets et bien localisés ;
le bilinéaire reste flou.

Spectres

EM4 $\approx$ eau (pic court, SAD 0.04). EM2 le plus dur (0.14).
EM1/EM3 resserrés sur la référence.

BILAN & PERSPECTIVES

Un upsampling sans label, fait pour le démixage non-supervisé

Conclusion

NAF entraîné sur DOFA gelé, par self-distillation MSE.

Des features grossières deviennent spatialement résolues, sans aucun label.

Perspectives

Entraîner NAF avec plus d'images hyperspectrales (ce qui demande plus de ressources)

Utiliser d'autres méthodes d'entraînement de NAF (par exemple crop & stitch)