

Comment les algorithmes de recommandation apprennent et utilisent les préférences politiques des utilisateurs ?

Tim Faverjon

tim.faverjon@sciencespo.fr

SciencesPo
MÉDIALAB

Antielite-Salience

High anti-elite

Left

Right

Low anti-elite

— Distribution of top users by learned-dimension-3 (l_3)

- top 50.00 % ($l_3 > 0.01$)
- top 25.00 % ($l_3 > 0.43$)
- top 12.50 % ($l_3 > 2.78$)
- top 6.25 % ($l_3 > 8.35$)
- top 3.12 % ($l_3 > 16.07$)
- top 1.56 % ($l_3 > 23.51$)
- top 0.78 % ($l_3 > 27.54$)

— Distribution of top users by learned-dimension-4 (l_4)

- top 50.00 % ($l_4 > 0.01$)
- top 25.00 % ($l_4 > 0.70$)
- top 12.50 % ($l_4 > 4.07$)
- top 6.25 % ($l_4 > 9.32$)
- top 3.12 % ($l_4 > 15.49$)
- top 1.56 % ($l_4 > 22.01$)
- top 0.78 % ($l_4 > 26.68$)

Digital Services Act

Demande aux plateformes

Chap. III, sec. 3, *art. 27*

1. Expliquer le fonctionnement des algorithmes de recommandations
(paramètres importants, raison des classement...)

Chap. III, sec. 3, *art. 26*

2. Ne pas discriminer selon certaines catégories de données personnelles
(en particulier les opinions politiques)

Digital Services Act

Demande aux plateformes

Chap. III, sec. 3, *art. 27*

1. Expliquer le fonctionnement des algorithmes de recommandations
(paramètres importants, raison des classement...)

Comment trouver les *paramètres importants* de l'algorithme ?

Chap. III, sec. 3, *art. 26*

2. Ne pas discriminer selon certaines catégories de données personnelles
(en particulier les opinions politiques)

Comment s'assurer les **opinions politiques** ne sont pas des *paramètres importants* ?

Comment comprendre les algorithmes de recommandations ?

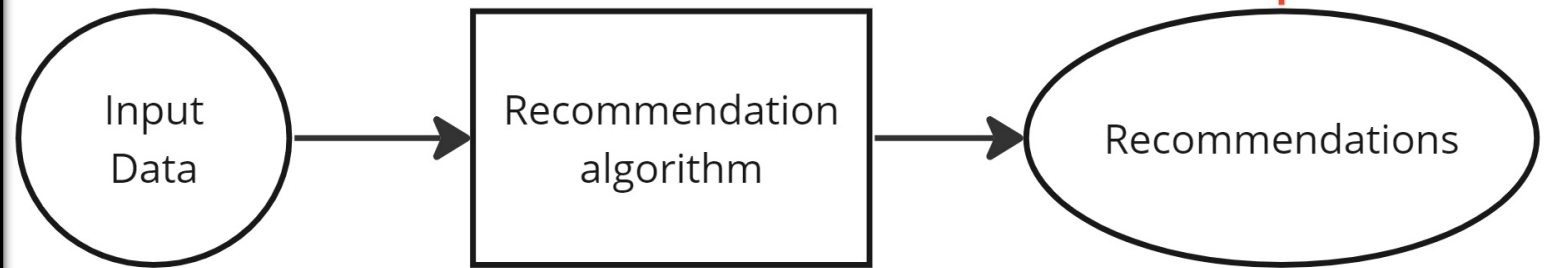
Deux familles de méthodes

Audit :

On regarde les *résultats* des recommandations.

Audit :

Are the recommendations as expected ?



Explication :

On explique le fonctionnement du *modèle* de recommandation.

Explication :

Which features impact the model and the recommendations ?

L'idée :

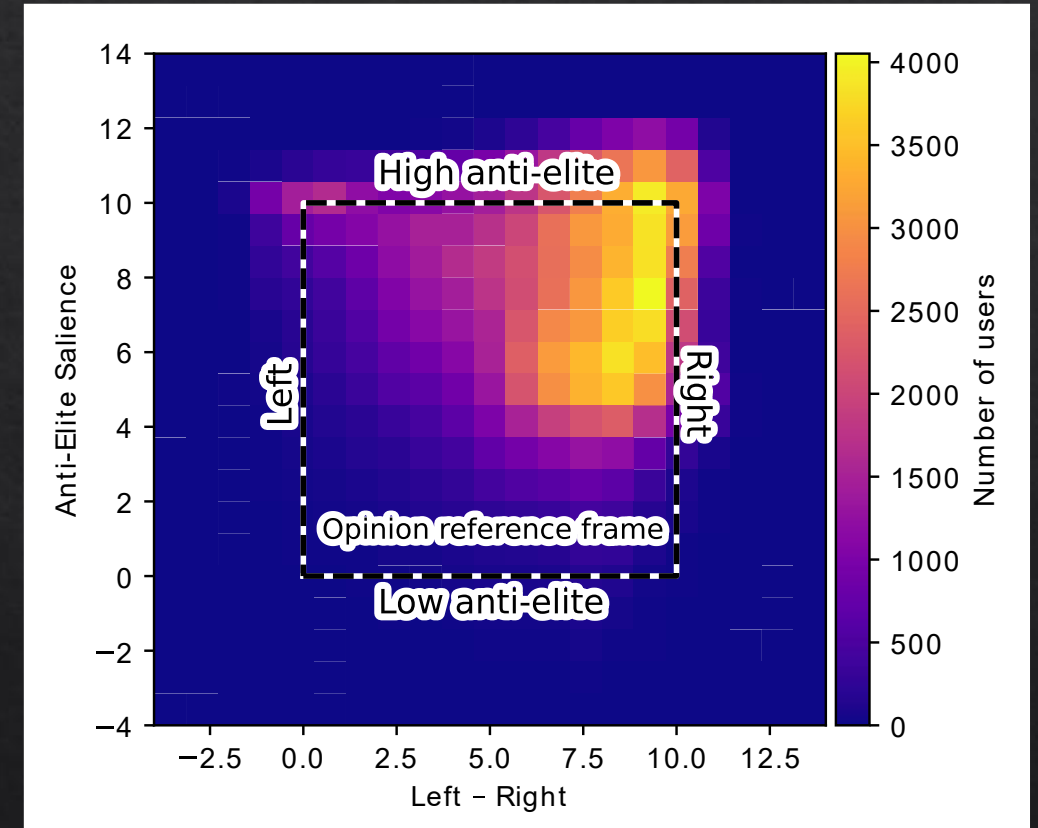
- ◆ Construire une méthode d'explication pour les algorithmes de recommandation
- ◆ Utilisant les opinions politiques des utilisateurs

Mesurer les attitudes politiques

- ◇ On construit un espace politique à partir du réseau d'utilisateurs suivant des parlementaires sur Twitter
- ◇ On donne un sens politique aux dimensions de cet espace en les reliant aux enquêtes d'experts (CHES)

Pour la France on obtient deux attitudes politiques qui expliquent le mieux l'espace politique :

1. **Left – Right** : droite et gauche classiques
2. **Anti-elite Salience** : défiance envers les institutions (0 = aucune défiance)



L'idée :

- ◇ Construire une méthode d'explication pour les algorithmes de recommandation
- ◇ Utilisant les opinions politiques des utilisateurs

La méthode :

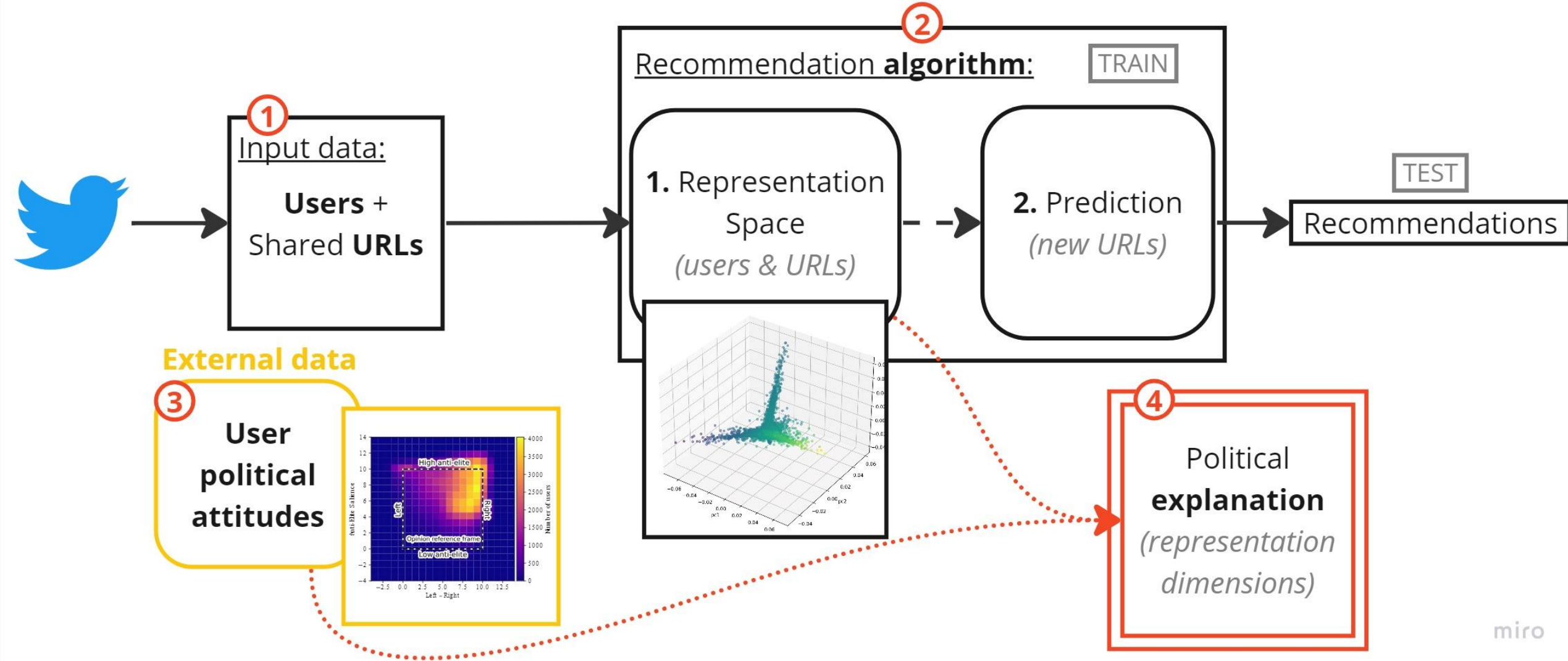
Une preuve de concept :

- ◇ Algorithme connu
- ◇ Données réelles (avec mesures politiques)

Le résultat :

Le modèle contient des information politiques sur les utilisateurs

La méthode : Preuve de concept



miro

Résultats

Certaines dimensions du modèle sont corrélées avec les attitudes politiques des utilisateurs.

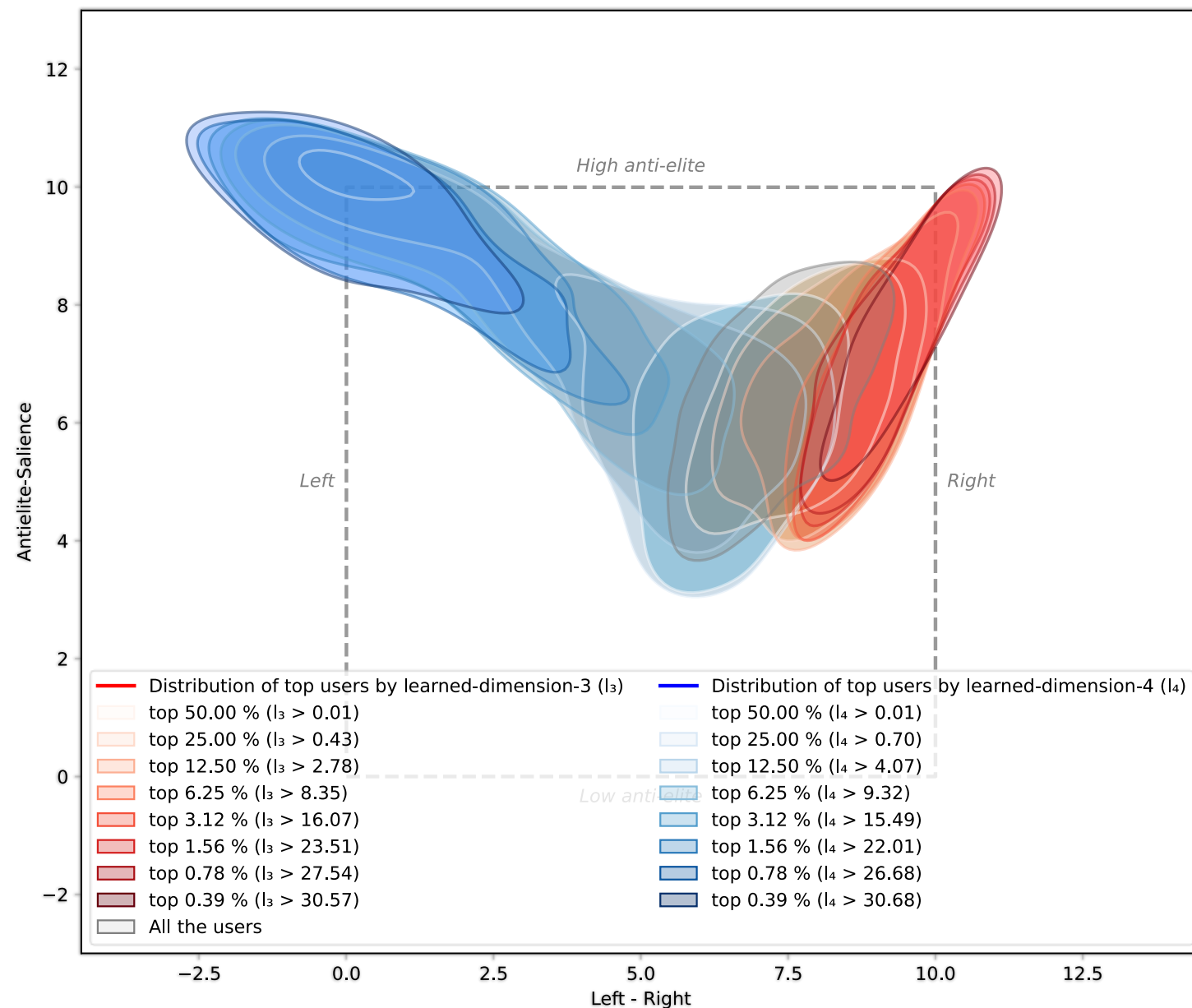
Exemple :

Le modèle identifie des groupes d'utilisateurs (**rouge** et **bleu**)

qui sont en réalité des groupes politisés (position)

◇ **gauche** – anti-elite

◇ **droite** – anti-elite



Résultats

- ◊ Un algorithme de recommandation peut apprendre et utiliser les opinions politiques des utilisateurs **sans que cela soit explicitement demandé**
- ◊ Nous pouvons mesurer ce phénomène

Implication de notre méthode :

- ◊ Identifier les contournements du DSA (discrimination politique) directement dans les modèles
- ◊ Envisager la construction d'outils et de solutions techniques

Comment les **algorithmes de recommandation** apprennent et utilisent les **préférences politiques** des utilisateurs ?

Tim Faverjon : tim.faverjon@sciencespo.fr

