

Identification de plusieurs fonctions de valeurs additives à partir de préférences exprimées anonymement

Vincent Auriou^{1,2}, Khaled Belahcène¹, Emmanuel Malherbe², Vincent Mousseau¹, Marc Pirlot³

¹ MICS, CentraleSupélec, Université Paris Saclay, France

² Artefact Research Center, France

³ MATHRO, Université de Mons, Belgique

vincent.auriau@artefact.com, khaled.belahcene@centralesupelec.fr, emmanuel.malherbe@artefact.com,
vincent.mousseau@centralesupelec.fr, marc.pirlot@umons.ac.be

Résumé

Eliciter un modèle de préférences additif consiste à poser une suite de questions à un décideur pour identifier une fonction de valeur additive représentant ses préférences. Dans ce travail, nous considérons un groupe de décideurs avec pour objectif d'éliciter une fonction de valeur pour représenter les préférences de chacun. Nous considérons le cas où l'on questionne deux décideurs, mais sans savoir quelle réponse correspond à quel décideur. Nous proposons une procédure d'élicitation identifiant les deux modèles de préférences lorsque les fonctions marginales représentées par des fonctions additives linéaires par morceaux.

Mots-clés

Modèle de préférence additif linéaire par morceaux, élicitation de modèles multiples, identification.

1 Présentation du problème et procédure d'élicitation

Le modèle de valeur additive est un choix standard pour spécifier des préférences sur un ensemble d'alternatives évaluées selon plusieurs critères, voir [3]. Plus précisément, étant donnée une alternative x définie par son évaluation sur n critères $x = (x_1, \dots, x_n)$, où x_i est l'évaluation de x sur le critère i , la valeur de x est définie par $u(x) = \sum_i u_i(x_i)$, où u_i est la fonction de valeur marginale sur le critère i . Dans cet article, nous considérons que ces fonctions marginales u_i sont strictement croissantes et linéaires par morceaux, avec une décomposition linéaire définie a priori. Dans la suite, nous appellerons un tel modèle de préférence un modèle UTA. Afin d'éliciter un modèle UTA, il est courant d'utiliser des requêtes standards correspondant à des questions de type "matching" [4], impliquant deux alternatives x et y dont les évaluations ne diffèrent que sur deux critères i et j . Trois des évaluations parmi x_i, x_j, y_i , et y_j sont fixées, et le décideur répondant fournit la valeur manquante de manière à ce que x et y soient indifférents (noté $x \sim y$). Cela implique que $u_i(x_i) + u_j(x_j) = u_i(y_i) + u_j(y_j)$. L'ensemble des points (x_i, x_j) dans le plan (i, j) pour lesquels $u_i(x_i) + u_j(x_j)$

est constant forme une courbe d'indifférence dans ce plan. De telles questions de matching, permettent identifier les points appartenant à une courbe d'indifférence [2], ce qui sera utile dans la suite.

Nous considérons un framework dans lequel l'objectif est d'éliciter simultanément deux modèles UTA. Dans ce contexte, la réponse à une requête de matching consiste en deux évaluations a et a' pour la composante non fixée. Ces réponses sont fournies de manière anonyme, sans possibilité de savoir à quel décideur appartient chaque réponse. Le but de cet article est d'étudier ce problème sous l'angle de l'identifiabilité. Autrement dit, nous cherchons à savoir s'il est toujours possible de définir une suite finie de requêtes dont les réponses permettent d'identifier les deux modèles de valeur additive linéaire par morceaux. Ce problème d'identifiabilité peut aussi être formulé comme un jeu dans lequel un premier joueur spécifie les requêtes et un second joueur fournit des réponses (véridiques). La question devient alors : existe-t-il une stratégie gagnante finie pour le premier joueur ? Une autre variante du jeu (non traitée ici) permettrait au second joueur de répondre avec un bruit (adversarial).

Bien que le problème soit formulé de manière purement formelle, il se retrouve bien dans des cas d'usages concrets. Par exemple, un supermarché souhaite généralement adapter la liste de ses produits à sa base de clientèle et sélectionner ces produits sur la base de leurs préférences. Toutefois, une telle population montre généralement des comportements et préférences très différents et il est standard de considérer une segmentation de la clientèle. Chaque segment représente alors un groupe de clients aux préférences homogènes, et dans notre cas un décideur dont la fonction de valeur est à éliciter.

Exemple 1 Dans cet article, nous utiliserons comme fil conducteur l'expression de préférences pour le choix d'une voiture électrique selon deux critères : l'autonomie (de 100 km à 600 km) et le prix (de 10 000€ à 50 000€). Les réponses aux requêtes de préférence sont obtenues via un système en ligne garantissant l'anonymat. Deux décideurs expriment leurs préférences : Commuter utilise la voiture

dans un contexte urbain pour de courtes distances, tandis que Traveler effectue de longs trajets. Commuter et Traveler disposent respectivement de 30 000€ et 40 000€ d'épargne liquide ; ils doivent contracter un prêt pour les voitures dont le prix dépasse leur épargne.

Ce type de problème a déjà été abordé dans une perspective d'apprentissage de modèle de préférence [1]. Le résultat d'identifiabilité proposé ici est fondamental pour justifier la recherche d'algorithmes efficaces permettant d'inférer plusieurs modèles UTA à partir de déclarations de préférences. Notre objectif est de fournir une procédure d'élaboration systématique permettant d'identifier deux fonctions de valeur additives linéaires par morceaux à partir d'une séquence finie de requêtes de matching. Dans cette introduction, nous proposons une présentation générale de cette procédure. Le principal résultat de l'article peut se formuler comme suit.

Theorem 1.1 *Il est possible d'élucider deux modèles de préférence additifs linéaires par morceaux à partir d'informations de préférence anonymisées en utilisant un nombre fini de requêtes d'indifférence.*

Dans notre cadre, chaque requête de matching, ou d'indifférence, implique des alternatives qui varient uniquement selon deux critères i et j . Trois valeurs sont fixées par l'Éliciteur, et les réponses sont fournies par les différents Décideurs de telle sorte à ce qu'elles correspondent à deux couples d'alternatives indifférentes. L'espace de valeur des critères (i, j) définit un plan où les requêtes peuvent être entièrement représentées, tout le reste étant égal. En particulier, étant donné que les fonctions de valeur marginales sont linéaires par morceaux, les courbes d'indifférence sont également représentées par des fonctions linéaires par morceaux (différentes). Dans ce plan des critères, il est utile de définir deux types de requêtes à la géométrie spécifique que nous utiliserons comme «blocs de construction» de notre procédure d'élaboration :

Requête sur un seul rectangle : les valeurs de la requête ainsi que les réponses sont contenues dans un seul rectangle défini par les segments linéaires, définis a priori, sur les deux échelles de critères.

Requête sur des rectangles adjacents : les valeurs de la requête sont choisies de manière à impliquer deux rectangles voisins définis par les segments linéaires sur les échelles de critères.

Les requêtes sur un seul rectangle permettent d'élucider les pentes anonymisées des fonctions de valeur marginales sur un intervalle donné, tandis que les requêtes sur des rectangles adjacents fournissent des contraintes de couplage permettant d'assigner les pentes aux deux fonctions de valeur.

En utilisant ces requêtes standards, la procédure peut être résumée comme suit :

Initialisation : la procédure commence avec les deux premiers critères et postule (sans perte de généralité) que la pente sur le critère 1 est égale à un sur le premier seg-

ment. Une requête sur un seul rectangle identifie la pente de chaque modèle pour le premier segment du critère 2.

Itération : étant donné un segment ℓ élaboré sur le critère 1, une succession de requêtes permet d'identifier les pentes des deux modèles sur le segment $\ell + 1$; il en va de même pour le critère 2.

Généralisation : pour obtenir l'ensemble des fonctions de valeur marginales, on applique le même principe à toutes les paires de critères $(1, i)$, $i \neq 1$.

La suite de l'article est organisée comme suit : après une revue des travaux connexes en Section 2, la Section 3 introduit les notations générales. La Section 4 présente les deux types de requêtes servant de blocs de construction pour la procédure d'identification. La Section 5 présente les résultats justifiant la nature itérative de la procédure générale, développée en Section 6.

Références

- [1] Vincent Auriau, Khaled Belahcène, Emmanuel Malherbe, and Vincent Mousseau. Learning multiple multicriteria additive models from heterogeneous preferences. In Rupert Freeman and Nicholas Mattei, editors, *Algorithmic Decision Theory*, pages 207–224. Springer, 2025.
- [2] Denis Bouyssou and Marc Pirlot. Conjoint measurement tools for mcdm. In Salvatore Greco, Matthias Ehrgott, and José Rui Figueira, editors, *Multiple Criteria Decision Analysis : State of the Art Surveys*, pages 97–151. Springer New York, 2016.
- [3] D.H. Krantz, R.D. Luce, P. Suppes, and A. Tversky. *Foundations of Measurement - Volume 1 : Additive and Polynomial Representations*. Academic Press, Incorporated, 1971.
- [4] D. von Winterfeldt and W. Edwards. *Decision analysis and behavioral research*. Cambridge University Press, Cambridge, 1986.