

Learning multiple multicriteria additive models from heterogenous preferences

Vincent Auriau^{†,*}, Emmanuel Malherbe^{*}, Vincent Mousseau[†]

[†] CentraleSupélec - Université Paris Saclay, ^{*} Artefact Research Center

Mots-clés : *Additive preference model, preference learning, heterogeneous learning set*

1 Introduction and motivation

A standard setting for preference learning [3] in multicriteria ranking problems with an additive representation of preferences dates back from the UTA method [2]. In this seminal work, an additive piece-wise linear model is inferred from a set of a learning set composed of pairwise comparisons. In this setting, the learning set is provided by a single Decision-Maker (DM) and UTA infers the additive model that best match the learning set, even when with noisy preferences. However, the learning set is supposed to be provided by a single DM. We extend this framework to the case in which (i) multiple DMs having each their own preferences provide part of the learning set, and (ii) the learning set is provided as a whole without the knowledge of which DM expressed each pairwise comparison. Hence, the problem amounts at inferring a preference model for each DM, which simultaneously requires to “discover” the segmentation of the learning set. Let us consider the two following illustrative examples.

Example 1 (clients preferences in a supermarket): a supermarket is willing to adapt the list of products to its customer base. The products selected should align with the client’s preferences. Yet, all customers do not necessarily have the same preferences, and it is standard to consider a market segmentation in which each segment represents a group of clients with homogeneous preferences. Each store can rely on actual sales to identify the segmentation and learn the clients’ preferences. For each sale of a product x , one can derive a comparison $x \succ y$ for each product y present in the store and substitutable to x . Hence, the problem can be formulated as clustering the clients and learning a preference model for each cluster.

Example 2 (Comparisons of films on social media): A website about cinema presents several films and request feedback from users. The users provide scores for films, from which pairwise comparisons can be derived. The users are not requested to have an account and are hence unidentifiable. For marketing purposes, the website wants to identify, from the collected data, k clusters of users with each a specific preference model.

2 Formal setting and notations

We consider a ranking problem with n criteria. The evaluation scale of criterion i , $i \in \{1..n\}$, is denoted X_i . Hence an alternative $x = (x_1, \dots, x_n)$ is an element of $\prod_i X_i$. We are provided with a learning set of P pairs $(x^{(j)}, y^{(j)})$ where $x^{(j)} = (x_1^{(j)}, \dots, x_n^{(j)})$ is preferred to $y^{(j)} = (y_1^{(j)}, \dots, y_n^{(j)})$, $j = 1..P$. We aim to represent this learning set using K additive piece-wise linear models (with L linear segments). Each of these additive models is defined by the marginal value of each breakpoint on each criterion, i.e., $u_i^k(x_i^l)$, $k \in \{1..K\}$, $i \in \{1..n\}$, $l \in \{1..L\}$. A pair $(x^{(j)}, y^{(j)})$ is considered correctly represented if it is the case for at least one of the K additive value models $u^k(\cdot)$, i.e. $u^k(x^{(j)}) > u^k(y^{(j)})$.

3 Algorithms for preference learning/segmentation

We propose a mathematical programming formulation of the problem that takes as input a learning set of learning set of P pairs $(x^{(j)}, y^{(j)})$, a number of clusters $K \in \mathbb{N}$, and a number of linear piece $L \in \mathbb{N}$, and returns K piecewise linear preference models (with L linear segment). The decision variables are the following:

- $u_i^k(x_i^l)$, $\forall k = 1..K$, $\forall i = 1..n$, $\forall l = 0..L$, defining the K UTA models,
- $z^k(j) \in \{0, 1\}$, $z^k(j) = 1$ if the pair $(x^{(j)}, y^{(j)})$ belongs to cluster k , else $z_k(j) = 0$, i.e.,
 $z^k(j) = 1 \Rightarrow \sum_{i=1}^n u_i^k(x_i^{(j)}) > \sum_{i=1}^n u_i^k(y_i^{(j)})$
- $\sigma^+(x^{(j)}), \sigma^-(x^{(j)}), \sigma^+(y^{(j)}), \sigma^-(y^{(j)}) \geq 0$, error variables relative to alternatives present in the learning set.

The constraints should enforce normalization [a] and monotonicity [b] of linear models, guaranty the correct definition of $z^k(j)$ [c], and ensure that at least on linear model represents each comparison of the learning set [d]:

- [a] $u_i^k(x_i^0) = 0, \forall i \in [1, n], \sum_{i=1}^n u_i^k(x_i^L) = 1$
- [b] $u_i^k(x_i^{l+1}) \geq u_i^k(x_i^l), \forall i \in [1, n], \forall l \in [0, L-1]$
- [c] $u^k(x^{(j)}) - u^k(y^{(j)}) + M(1 - z^k(j)) \geq 0, j \in \{1, \dots, P\}$
- [d] $\sum_{k=1}^K z^k(j) \geq 1, j \in \{1, \dots, P\}$

In this mathematical programming formulation, the objective is to minimize errors, i.e., $\min \sum_j (\sigma^+(x^{(j)}) + \sigma^-(x^{(j)}) + \sigma^-(y^{(j)}) + \sigma^+(y^{(j)}))$. The resolution of the above mathematical program using a solver can be computationally prohibitive. Therefore we have designed a heuristic algorithm. Inspired by Expectation-Maximization [1], it works with the repetition of two successive steps after a random initialization of K UTA models:

1. Assign each comparison $(x^{(j)}, y^{(j)})$ to the additive model with the highest difference $u^k(x^{(j)}) - u^k(y^{(j)})$
2. Infer the K UTA models for each cluster independently, using its corresponding set of comparisons

The steps are repeated until either errors are null or the partition of the learning set does not change after an iteration.

4 Numerical results

In the presentation, we will present the results of numerical experiments to solve synthetic datasets using both the mathematical programming formulation and the heuristic algorithm.

5 Conclusion

This work presents a new preference learning/segmentation problem with multiple application domains. We proposed a mathematical programming resolution scheme and a heuristic algorithm whose performance have been tested through extensive numerical experiments.

References

- [1] A. Dempster, N. Laird, and D. Rubin. Maximum likelihood from incomplete data via the EM algorithm. *Journal of the Royal Statistical Society: Series B (Methodological)*, 39(1):1–22, 1977.
- [2] E. Jacquet-Lagrange and J. Siskos. Assessing a set of additive utility functions for multicriteria decision-making, the UTA method. *EJOR*, 10(2):151–164, 1982.
- [3] V. Mousseau and M. Pirlot. Preference elicitation and learning. *EURO Journal on Decision Processes*, 3(1-3), May 2015.