

STATISTIQUE ET ANALYSE DES DONNÉES

E. JACQUET-LAGREZE

Méthodes explicatives en analyses de préférences ordinales

Statistique et analyse des données, tome 2, n° 2 (1977), p. 45-58

http://www.numdam.org/item?id=SAD_1977__2_2_45_0

© Association pour la statistique et ses utilisations, 1977, tous droits réservés.

L'accès aux archives de la revue « Statistique et analyse des données » implique l'accord avec les conditions générales d'utilisation (<http://www.numdam.org/conditions>). Toute utilisation commerciale ou impression systématique est constitutive d'une infraction pénale. Toute copie ou impression de ce fichier doit contenir la présente mention de copyright.

NUMDAM

Article numérisé dans le cadre du programme
Numérisation de documents anciens mathématiques
<http://www.numdam.org/>

E. JACQUET-LAGREZE

Université Paris IX Dauphine.

0 - INTRODUCTION

- La plupart des modèles développés en analyse multicritère, l'ont été dans la perspective d'être des outils d'aide à la décision. L'analyse de la préférence d'un décideur se fait en déterminant une famille cohérente de critères (cf. Roy (1975-1)) et en construisant un modèle de la préférence globale agrégeant les critères : construction d'une fonction de valeur (cas certain) d'une fonction d'utilité (cas aléatoire : prise en compte du risque) (cf. Von Winterfeldt and Fisher (1975)), d'une relation de surclassement (cf. Roy (1975), Jacquet-Lagrèze (1975)), mise en place d'une méthode interactive de recherche d'un compromis (cf. Zeleny (1974),...). Cette première approche dans la modélisation correspond à une attitude normative ou prescriptive.

- Confronté au problème de la description d'un processus de décision ou même plus simplement à celui de trouver un modèle expliquant le comportement d'un individu en situation de choix, l'homme d'étude doit adopter une attitude plus descriptive. Dans un problème de choix mettant en oeuvre plusieurs critères, le problème peut d'abord s'énoncer comme : "Essayez de comprendre pourquoi telle action a été choisie" ou plus précisément : "Comment expliquer une préférence globale spontanée à l'aide d'une famille de critères donnée et d'une règle d'agrégation de ces critères donnée". Ce problème apparaît comme le problème inverse de l'agrégation des critères : c'est un problème de désagrégation. Cette seconde attitude dans la modélisation a été celle de psychologues (cf. Slovic & Lichtenstein (1971)) et a fait l'objet d'un grand nombre d'applications expérimentales (cf. Bowman (1963), Fischer (1975), Huber et al. (1969), 1961),...). La plupart des travaux réalisés utilisent des modèles d'analyse statistique tel que la régression linéaire multiple :

Soient A un ensemble d'actions (objets à comparer, décisions envisagées,...) $\underline{g}(a) = \{g_1(a), \dots, g_n(a)\}$ une famille cohérente de critères ($a \in A$).

* - Cette recherche et la mise au point du programme ADEPREF ont été réalisées dans le cadre du projet COREF/DGRST. N° 75-7-0230.

On suppose que le modèle de la préférence globale est une fonction de valeur additive $V(a) = \sum_{i=1}^n p_i f_i [g_i(a)]$ où p_i est le taux de substitution entre le critère transformé i et un critère de référence et f_i une transformation monotone des critères g_i . On suppose également que la préférence spontanée observée est une variable quantitative $y(a)$, alors le jeu de poids $\{p_i\}$ peut être estimé en utilisant la régression linéaire multiple de façon à ce que $y(a)$ soit le plus proche possible de $V(a)$ au sens des moindres carrés.

- Cependant dans bien des cas, dès qu'il s'agit de préférences, la préférence spontanée ne peut être appréhendée qu'au moyen d'une variable ordinale et plus généralement au moyen d'une relation binaire qui peut présenter des situations de non comparaison et même d'intransitivité. C'est l'objet de cette recherche que de présenter quelques modèles explicatifs d'une relation binaire. Dans la section 1, on présente un des modèles utilisant une fonction de valeur additive et dans la section 2 un modèle basé sur une agrégation ordinale des critères. La section 3 est une courte discussion sur ce problème.

1. - MODELES EXPLICATIFS BASES SUR DES FONCTIONS DE VALEUR

Comme dans le modèle de la régression linéaire multiple, on suppose dans cette section que le modèle explicatif est une fonction de valeur additive $V(a) = \sum_{i=1}^n p_i f_i [g_i(a)]$. On suppose par contre que la préférence spontanée à expliquer n'est connue que par une relation binaire $R = (P, I)$ définie sur $A \times A$ où P , la partie antisymétrique de R représente la préférence stricte, et I la partie symétrique de R représente l'indifférence. Si R est une relation de pré-ordre total alors on est dans le cas particulier du problème de l'explication d'une variable ordinale. Dans les deux modèles présentés ci-dessous, le problème consiste à trouver un jeu de poids $p = \{p_1, \dots, p_n\}$ tel que $V(a)$ soit aussi "compatible" que possible avec la relation R .

1.1. Le modèle ORDREG (*) de Srinivasan et Shocker (1973)

Dans ce modèle les auteurs considèrent une simple forme additive $V(a) = \sum_i p_i g_i(a)$. La compatibilité entre $V(a)$ et R est définie à l'aide de la condition (1) :

$$(1) \quad (a_k, a_\ell) \in P \implies V(a_k) \geq V(a_\ell)$$

Les auteurs supposent qu'il est en général impossible de trouver un jeu de poids $\{p_i\}$ satisfaisant à toutes les conditions du type (1). La méthode consiste alors à trouver un jeu de poids minimisant une fonction d'erreurs à l'aide du programme linéaire suivant :

$$\left\{ \begin{array}{l} \text{Min} = \sum_{(a_k, a_\ell) \in P} F z_{k\ell} \\ (2) \quad \sum_{i=1}^n p_i [g_i(a_k) - g_i(a_\ell)] + z_{k\ell} \geq 0 \text{ pour tout } (a_k, a_\ell) \in P \\ (3) \quad \sum_{i=1}^n p_i \sum_{(a_k, a_\ell) \in P} [g_i(a_k) - g_i(a_\ell)] = 1 \\ p_i \geq 0, \quad z_{k\ell} \geq 0. \end{array} \right.$$

Les contraintes (2) introduisent des variables d'erreur $z_{k\ell}$ pour celles des contraintes résultant de la condition (1) qui ne sont pas satisfaites : $z_{k\ell} > 0 \implies V(a_k) - V(a_\ell) < 0$ alors que $(a_k, a_\ell) \in P$.

La contrainte (3) évite que l'on obtienne la solution $p_i = 0$ pour tout i .

Ayant utilisé ce modèle, je voudrais présenter quelques difficultés que l'on peut rencontrer.

1). En utilisant une somme pondérée des g_i , il faut supposer que ces derniers sont des graduations c'est-à-dire que si $g_i(a) - g_i(b) = g_i(c) - g_i(d)$, la supériorité de a sur b doit être

(*) ORDinal REGression, modèle inclus dans le programme LINMAP.

jugée la même que celle de c sur d . Il est alors important d'introduire éventuellement des transformations monotones f_i des critères g_i .

2). Si le nombre des critères est élevé comparé au nombre d'actions comparées $|A|$, il peut exister une infinité de solutions $\{p_i\}$ donnant toutes les valeurs $F=0$, certaines d'entre elles pouvant être très contrastées. Ceci est important puisqu'on s'intéresse dans la suite de l'analyse à ces solutions p_i et non pas à F .

3). Seule la partie antisymétrique de R est prise en compte. L'indifférence I ne l'est pas.

4). La contrainte (3) paraît arbitraire. Je pense que dans la plupart des applications, on peut utiliser une contrainte ayant plus de sens.

Néanmoins ce modèle est extrêmement intéressant, et on peut aisément éviter ces difficultés en lui apportant quelques modifications.

1.2. Quelques extensions au modèle ORDREG

- Tout d'abord il faut pouvoir introduire des transformations monotones f_i de g_i de façon à obtenir des graduations (*)

- Pour éviter la contrainte (3), je propose la solution suivante. Supposons que $f_i(g_i)$ soit une graduation, alors toute transformation linéaire de f_i est aussi une graduation. Considérons alors la transformation particulière :

$$v_i(a) = \frac{f_i[g_i(a)] - f_{i*}}{f_i^* - f_{i*}}$$

où $f_i^* = \max_{a \in A} f_i[g_i(a)]$, $f_{i*} = \min_{a \in A} f_i[g_i(a)]$

(+) une autre modification, plus importante celle-là, consisterait à trouver la meilleure transformation monotone directement, en sortie du modèle

En sortie d'un tel modèle, on peut présenter entre le minimum de F et le jeu de poids $\{p_i\}$ cherché, la relation de préordre $R' = (P', I')$ correspondant à la fonction de valeur $V(a)$:

$$V(a_k) > V(a_\ell) \iff (a_k, a_\ell) \in P'$$

$$V(a_k) = V(a_\ell) \iff (a_k, a_\ell) \in I'$$

Il est alors possible de comparer la relation initiale R et le préordre R' ainsi obtenu, à l'aide de la différence symétrique, ou mieux, avec les nombres en pourcentages de transformations des couples $A \times A$.

$n(P, P')$: pourcentage de préférences inchangées

$n(I, I')$: pourcentage d'indifférences inchangées

$k(P, P')$: pourcentage de préférences strictes inversées

$n(P, I')$: pourcentage de préférences de R changées en indifférences dans R'

$n(I, P')$: pourcentage d'indifférences de R' changées en préférences dans R

$n(R-PUI, P')$: pourcentage de paires non comparées transformées en préférence stricte dans R'

$n(R-PUI, I')$: pourcentage de paires non comparées transformées en indifférences dans R' .

2.- UN MODELE EXPLICATIF UTILISANT UNE PROCEDURE D'AGREGATION ORDINALE DES CRITERES.

Dans cette section, le modèle de la préférence spontanée à expliquer est encore une relation binaire $R=(P,I)$ totale ou partielle. Mais on suppose maintenant que les critères explicatifs sont ordinaux, ou bien, s'ils ont une structure plus riche, on ne s'intéresse qu'aux préordres qu'ils définissent sur A . On suppose plus généralement que les variables explicatives de R sont n relations binaires $\{r^1, \dots, r^i, \dots, r^n\}$.

On suppose que ces relations sont probabilistes :

$$r_{k\ell}^i + r_{\ell k}^i = 1 \quad \forall (a_k, a_\ell) \in A \times A$$

Dans des relations probabilistes, on représente les situations de préférence et d'indifférence de la façon suivante :

$$r_{k\ell}^i = 1, r_{\ell k}^i = 0 \iff a_k \succ_i a_\ell$$

$$r_{k\ell}^i = r_{\ell k}^i = 1/2 \iff a_k \sim_i a_\ell$$

Une situation de préférence large serait représentée par :

$$1/2 < r_k^i < 1$$

2.1. La procédure d'agrégation

On peut retenir une classe de procédure d'agrégation à seuil généralisant la procédure majoritaire de Condorcet et la règle de l'unanimité de Pareto (cf. Jacquet-Lagrèze (1973)).

Soit $\{p_1, \dots, p_i, \dots, p_n\}$ un jeu de poids normé :

$$\sum_{i=1}^n p_i = 1$$

Soit g la relation barycentrique définie par

$$g_{k\ell} = \sum_{i=1}^n p_i r_{k\ell}^i, \quad g \text{ est probabiliste.}$$

Considérons un seuil α ($1/2 \leq \alpha < 1$) et la procédure d'agrégation définie par

$$(6) \quad (a_k, a_\ell) \in P_\alpha \iff g_{k\ell} > \alpha \iff g_{\ell k} < 1-\alpha$$

$$(7) \quad (a_k, a_\ell) \in I_\alpha \iff 1-\alpha \leq g_{k\ell} \leq \alpha$$

Cette procédure donne une relation binaire $R_\alpha = (P_\alpha, I_\alpha)$ totale, mais non nécessairement transitive.

Pour $\alpha = 1/2$, cette méthode est la procédure majoritaire.

Pour $\alpha = 1-\epsilon$ (ϵ très petit), cette méthode est la règle de l'unanimité définie sur la préférence stricte.

2.2. Le modèle explicatif correspondant

Soit $R = (P, I)$ une relation binaire à expliquer à l'aide des n relations binaires $\{r^1, \dots, r^n\}$ et de la règle d'agrégation donnée ci-dessus.

Un jeu de poids $p = \{p_1, \dots, p_n\}$ est une solution exacte du problème si les conditions suivantes sont satisfaites :

$$(8) \quad \left\{ \begin{array}{l} (a_k, a_\ell) \in P \implies g_{k\ell} > \alpha \iff g_{\ell k} < 1-\alpha \\ (a_k, a_\ell) \in I \implies 1-\alpha \leq g_{k\ell} \leq \alpha \iff g_{k\ell} \leq \alpha \\ \text{et } g_{\ell k} \leq \alpha \end{array} \right.$$

Deux cas peuvent se présenter suivant qu'il existe ou non une solution exacte p (*)

1er cas : Il n'existe pas de solution p satisfaisant aux conditions (8)

Considérons les variables d'écart définies par :

$$(9) \quad \left\{ \begin{array}{l} g_{k\ell} - \alpha + z_{k\ell} \geq 0 \text{ pour tout } (a_k, a_\ell) \in P \\ g_{k\ell} - \alpha - z_{k\ell} \leq 0 \text{ pour tout } (a_k, a_\ell) \in I \end{array} \right.$$

Remarquons que I est symétrique et que $(a_k, a_\ell) \in I \implies (a_\ell, a_k) \in I \implies g_{\ell k} - \alpha - z_{\ell k} \leq 0$.

Le nombre de variables introduites par (9) est alors $|P| + |I|$ avec $m(m-1)/2 \leq |P| + |I| \leq m(m-1)$ si $m = |A|$ et si $R = (P, I)$ est totale.

Soit la fonction d'erreur définie par :

$$F = a \sum_{(a_k, a_\ell) \in P} z_{k\ell} + b \sum_{(a_k, a_\ell) \in I} z_{k\ell}$$

(*) On pourrait procéder de la même façon dans le modèle ORDREG.

Un vecteur $\underline{p}$ est solution du programme linéaire suivant :

$$\left\{ \begin{array}{l} \text{Min } F = a \sum_{(a_k, a_\ell) \in P} z_{k\ell} + b \sum_{(a_k, a_\ell) \in I} z_{k\ell} \\ \\ \sum_{i=1}^n p_i r_{k\ell}^i - \alpha + z_{k\ell} \geq 0 \text{ pour tout } (a_k, a_\ell) \in P \\ \\ \sum_{i=1}^n p_i r_{k\ell}^i - \alpha - z_{k\ell} \leq 0 \text{ pour tout } (a_k, a_\ell) \in I \\ \\ \sum_{i=1}^n p_i = 1 \\ \\ p_i \geq 0, z_{k\ell} \geq 0 \end{array} \right.$$

Comme dans le paragraphe 1.2., on peut choisir a et b de façon à traduire l'importance relative qu'on accorde dans le respect de la préférence et de l'indifférence dans les conditions (9). On peut également comparer la relation initiale R et la relation R' obtenue à l'aide du jeu de poids $\underline{p}$, des relations explicatives et de la règle d'agrégation donnée.

2ème cas : Il existe au moins une solution exacte $\underline{p}$ satisfaisant aux conditions (8)

Ce cas se produit lorsqu'on obtient dans le programme linéaire précédent la valeur 0 pour F .

Il est alors des plus important de caractériser au mieux l'ensemble des solutions admissibles pour $\underline{p}$, puisque c'est sur les valeurs comparées des p_i que porte ensuite l'analyse explicative.

L'ensemble des conditions (8) définit alors dans R^n un polyèdre $\mathcal{P}$ non vide qu'on peut chercher à explorer en maximisant et minimisant certaines formes linéaires.

$$\left\{ \begin{array}{l} \sum_{i=1}^n p_i r_{k\ell}^i > \alpha \quad \text{pour tout } (a_k, a_\ell) \in P \\ \sum_{i=1}^n p_i r_{k\ell}^i \leq \alpha \quad \text{pour tout } (a_k, a_\ell) \in I \\ \text{Max } \sum_{i=1}^n c_i p_i \quad \text{puis} \quad \text{Min } \sum_{i=1}^n c_i p_i \end{array} \right.$$

Si on considère m formes linéaires, on mettra ainsi en évidence $2m$ sommets de $\mathcal{P}$.

Si le n est faible, on pourra chercher les $2m$ sommets correspondant à $\text{Max } p_i$, $\text{Min } p_i$ pour $i=1, n$. Si n est plus élevé, on pourra se contenter d'un nombre plus faible de formes linéaires simples. On peut par exemple réaliser une classification des variables explicatives, retenir m classes (groupes) et choisir pour les m formes linéaires, la somme des poids des variables appartenant à un même groupe (variables plus fortement corrélées).

En suivant cette procédure, on peut caractériser l'ensemble des jeux de poids admissibles par $2m$ jeux de poids et le jeu de poids moyen $\underline{p}^m$ associé.

Le programme ADÉPREF^(*)

Il correspond au modèle explicatif présenté dans cette section. On a avantage à résoudre le problème dual en utilisant un code linéaire à variables bornées, la taille de la matrice des contraintes étant alors fortement réduite. (cf. Berges J.C., Jacquet-Lagrèze E. (1976)). Il a été utilisé jusqu'à maintenant sur un problème de 20 objets et 7 critères explicatifs.

3.- DISCUSSION ET EXEMPLE

Les modèles présentés ci-dessus appartiennent à la classe des modèles explicatifs. Une relation binaire est expliquée soit à l'aide d'une fonction de valeur additive, soit à l'aide d'un modèle d'agrégation

(*) Analyse par DESagrégation d'une PREFérence globale

de relations d'ordre, de relations binaires plus généralement.

Il serait intéressant de comparer ces modèles avec d'autres tels que MONANOVA (cf. Kruskal (1965)) ou MORALS (cf. Young et al.).

D'autres modèles analogues peuvent être construits en utilisant d'autres procédures d'agrégation notamment ordinales. Il serait également intéressant d'avoir des modèles prenant en compte le risque (désagrégation de fonctions d'utilités par exemple).

Enfin dans l'introduction, on a présenté ces modèles dans un contexte d'une analyse multicritère descriptive, voire explicative. Mais de tels modèles peuvent également être utilisés dans une approche normative ou prescriptive lorsqu'il s'agit par exemple de construire des critères en sous-agrégeant un groupe de critères, et cela lorsqu'il semble plus facile d'obtenir une préférence spontanée que des taux de substitution ou autres paramètres d'une procédure d'agrégation.

Pour illustrer ce dernier point considérons l'exemple suivant. Si on souhaite construire un critère de bruit dans une étude d'investissement routier, on peut supposer (suite à des études préalables) que cette nuisance est essentiellement fonction du trafic T et de distance d sans protection à l'axe routier. On peut chercher à agréger ces deux dimensions en une seule $g(T,d)$.

Pour réaliser une telle agrégation, on peut utiliser le modèle du paragraphe 1.2.

Soit A un ensemble de couples (T,d) : $a_1 = (T_1, d_1), \dots, a_k = (T_k, d_k)$

correspondant à différents niveaux du trafic et de la distance. On peut demander à un ensemble de personnes de comparer en terme de nuisance les différents couples a de A .

Après une telle étude expérimentale, il serait possible d'obtenir une relation binaire R sur $A \times A$ représentant le jugement spontané sur A . On peut alors essayer différentes fonctions d'agrégation telles que :

$$f(a) = f(T,d) = T e^{-\alpha d}$$

Une forme additive serait alors :

$$V(a) = \text{Log } [f(a)] = \text{Log } T - \alpha d$$

Le modèle du paragraphe 1.2. adapté à cette fonction de valeur serait ici :

$$\left\{ \begin{array}{l} \text{Min } F = a \sum_{(a_k, a_\ell) \in P} z_{k\ell} + b \sum_{(a_k, a_\ell) \in I} (z_{k\ell} + z_{\ell k}) \\ V(a_k) - V(a_\ell) + z_{\ell k} - z_{k\ell} = 0 \text{ pour tout } (a_k, a_\ell) \in R \\ z_{k\ell} \geq 0. \end{array} \right.$$

Ayant estimé le paramètre α de ce modèle, un critère g pourrait être construit sur cette échelle de nuisance en considérant la distribution $N^a(T, d)$ sur une telle échelle pour un projet routier a et un équivalent ponctuel tel que :

$$g(a) = \sum N^a(T, d) T e^{-\alpha d}.$$

REFERENCES

- BERGES J.C. & JACQUET-LAGREZE (1977) :
ADREPREF: "une méthode de désagrégation de préférences
ordinales", Note de travail COREF
- BOWMAN E.H. (1963) : "Consistency and Optimality in Managerial
Decision making" - Management Science 9.
- FISCHER G.W. (1975) : "Experimental applications of Multi-attribute
Utility Models" in Utility, Probability and Human Decision
Making, Edited by C. Wendt and C. Vlek, Reidel Publis-
hing company.
- HUBER G.P., DANESHGAR R., FORD D.L. (1971) : "An empirical comparison
of Five Utility Models for predicting jobs Preference".
Organizational Behavior and Human Performance 6.
- HUBER G.P., SAHNEY V.K., FORD D.L. (1969) : "A study of subjective
Evaluation Models" - Behavioral Science 14.
- JACQUET-LAGREZE E. (1973) : "Le problème de l'agrégation des préférences :
une classe de procédure à seuil", Mathématiques et
Sciences Humaines, 43.
- JACQUET-LAGREZE E. (1975) : "La modélisation des préférences, pré-ordres,
quasi-ordres et relations floues" - Thèse - Université
Paris V.
- KRUSKAL J.B. (1965) : "Analysis of factorial experiments by estimating
monotone transformation of the data" - Journal of the
Royal Statistical Society - Series B, 27
- ROY B. (1975-1) : From optimization to multicriteria decision aid : a
conceptual framework. Proceedings of the Jouy-en-Josas
Conference on Multiple Criteria Decision Making.
Edited by Thiriez and Zionts, Springer-Verlag, to appear.
- ROY B. (1975-2) : "Partial Preference analysis and decision aid : the
fuzzy outranking relation concept". Presented to the
IIASA workshop : Decision Making with Multiple Conflicting
Objective - Vienna, October 20-24.
- SLOVIC P., LICHTENSTEIN S. (1971) : "A comparison of Bayesian and Regression
Approaches to the study of Information Processing in
judgement " - Organizational Behavior and Human
Performance, 6.
- SRINIVASAN V., SHOCKER A.D. (1973) : "Estimating the weights for multiple
attributes in a composite criteria using pairwise
judgements" - Psychometrika, 38

- Von WINTERFELDT D., FISCHER G.W. (1975) : "Multi-attribute Utility theory : Models and assessment procedures" - in Utility, Probability and Human Decision Making, Edited by D. Wendt and C. Vlek, Reidel Publishing Company.
- YOUNG F.W., de LEEUW J., TAKANE Y. : "Multiple and Canonical regression with a mix of qualitative and quantitative variables" - Report n° 146, University of North Carolina.
- ZELENY M. (1974) : Linear multiobjective programming - Springer-Verlag.