

REVUE DE STATISTIQUE APPLIQUÉE

PH. COURCOUX

M. SÉMÉNOU

Une méthode de segmentation pour l'analyse de données issues de comparaisons par paires

Revue de statistique appliquée, tome 45, n° 2 (1997), p. 59-69

http://www.numdam.org/item?id=RSA_1997__45_2_59_0

© Société française de statistique, 1997, tous droits réservés.

L'accès aux archives de la revue « Revue de statistique appliquée » (<http://www.sfds.asso.fr/publicat/rsa.htm>) 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/>

UNE MÉTHODE DE SEGMENTATION POUR L'ANALYSE DE DONNÉES ISSUES DE COMPARAISONS PAR PAIRES

Ph. Courcoux, M. Séménou

ENITIAA/INRA, Unité de Statistique Appliquée à la Caractérisation des Aliments,
Rue de la Géraudière, BP 82225, 44322 Nantes Cedex 03

RÉSUMÉ

Les données de comparaisons par paires sont en général analysées en utilisant les modèles de Thurstone ou Bradley. Lorsque l'on suppose l'existence de plusieurs segments dans le panel de juges, on est amené à évaluer les configurations des stimuli dans chaque classe ainsi que les probabilités d'appartenance des individus aux différents segments. La méthode de classification proposée repose sur le modèle de Bradley. On utilise une technique d'estimation basée sur l'algorithme E.M. et une procédure de test de type Monte Carlo est mise en œuvre pour déterminer le nombre de classes homogènes. Cette méthode est appliquée à une étude de formulation d'un produit alimentaire de type cocktail.

Mots-clés : Comparaisons par paires, Modèle de Bradley, Classification, Algorithme E.M., Simulations de Monte Carlo.

ABSTRACT

Paired comparisons data are generally analyzed using Thurstone or Bradley models. The conventional approach is to consider the consumer group as an homogenous one and to estimate a single set of scale values of the stimuli. In order to classify the panel, a latent class segmentation method is presented, based on the Bradley-Terry-Luce model and using an E.M. algorithm. A Monte Carlo significance testing procedure can help to the determination of the number of homogenous classes.

Keywords : Paired comparisons, Bradley Model, Classification, E.M. Algorithm, Monte Carlo Simulations.

Introduction

Les épreuves de type comparaisons par paires sont employées en étude de marketing ou en évaluation sensorielle. Ces épreuves sont particulièrement adaptées quand on souhaite évaluer des produits de façon subjective. Un individu est alors soumis successivement à des paires de stimuli (produits, descripteurs,...) et doit choisir un élément de chaque paire selon le critère que l'on souhaite étudier. D'un point de

vue pratique, ces épreuves sont aisées à mettre en œuvre et fournissent une bonne qualité de discrimination (O'Mahony *et al*, 1994).

Pour n produits comparés par H consommateurs, le nombre total de comparaisons par paires est $\frac{Hn(n-1)}{2}$; lorsque le nombre de produits à évaluer devient important, l'étude exhaustive de ces paires est difficile et l'utilisation de plans de présentation incomplets s'avère nécessaire. Divers types de plans peuvent être employés parmi lesquels les plans incomplets équilibrés et les plans cycliques qui assurent l'équilibre des effets d'ordre et de report (David, 1988).

Le traitement des données de comparaisons par paires n'est pas un problème récent. Thurstone (1927) et Bradley-Terry-Luce (1952) ont développé des modèles couramment utilisés sous l'hypothèse d'homogénéité du panel de consommateurs et permettant de représenter ces données de manière unidimensionnelle. Plus récemment, des modèles de positionnement multidimensionnel ont été adaptés à ce type d'épreuves (Coombs, 1958; Schonemann et Wang, 1972; Carroll, 1980). Cette approche est employée dans le domaine du marketing et permet de rendre compte de la nature multidimensionnelle des produits comparés et de l'aspect individuel des jugements de préférence (De Soete et Carroll, 1983; De Sarbo *et al*, 1987).

Le problème de la classification des juges à l'issue d'épreuves de comparaisons par paires est important pour l'interprétation des données. L'utilisateur est en effet intéressé par la segmentation du panel en groupes homogènes. Cette méthode a été abordée en utilisant le modèle de Thurstone par De Soete et DeSarbo (1991) pour l'interprétation de données issues d'épreuves de type 1 parmi n ou par Dillon *et al* (1993) pour prendre en compte des caractéristiques des juges. Nous proposons une méthode basée sur le modèle de Bradley permettant d'étudier à la fois la segmentation du panel et les configurations des produits dans les différentes classes.

Présentation générale

Supposons que n objets (produits, descripteurs,...) soient présentés à H individus lors d'épreuves de comparaisons par paires. Chacun des H sujets doit choisir entre deux stimuli; la réponse de l'individu h pour la paire (i, j) est notée :

$$y_{ij,h} = \begin{cases} 0 & \text{si } j \text{ est préféré à } i \\ 1 & \text{si } i \text{ est préféré à } j. \end{cases}$$

$y_{ij,h}$ est l'observation d'une variable aléatoire $Y_{ij,h}$ suivant une loi de Bernoulli.

Soit $\pi_{ij,h}$ la probabilité que le consommateur h préfère i à j . En utilisant le modèle de Bradley-Terry-Luce (Bradley, 1952), cette probabilité devient :

$$\pi_{ij,h} = \frac{\pi_i}{\pi_i + \pi_j} \quad (1)$$

sous la contrainte : $\pi_i \in [0; 1]$ pour $i = 1$ à n et $\sum_{i=1}^n \pi_i = 1$.

Cela revient à exprimer les $\pi_{ij,h}$ en fonction des valeurs π_i (scores de Bradley) prises pour chacun des n objets. En supposant l'indépendance entre les variables $Y_{ij,h}$, la vraisemblance associée à l'ensemble des observations $y = {}^t(y_1, \dots, y_H)$ où $y_h = {}^t(y_{ij,h}/1 \leq i < j \leq n)$ est le vecteur d'observations pour l'individu h , s'écrit :

$$L(y; \pi) = \prod_{h=1}^H f(y_h, \pi) \quad (2)$$

où $\pi = {}^t(\pi_1, \dots, \pi_n)$ est le vecteur $1 \times n$ des paramètres inconnus et $f(y_h, \pi)$ est défini par :

$$f(y_h, \pi) = \prod_{i=1}^{n-1} \prod_{j=i+1}^n \left(\frac{\pi_i}{\pi_j} \right)^{y_{ij,h}} \left(\frac{\pi_j}{\pi_i + \pi_j} \right). \quad (3)$$

L'estimation de π est obtenue par la méthode du maximum de vraisemblance (Dykstra, 1956).

Cette démarche suppose *a priori* l'homogénéité du panel, ce qui, dans certains cas, peut être critiquable. Nous proposons une méthode permettant non seulement de tenir compte d'une éventuelle hétérogénéité du jury, mais également, après une étape de segmentation, de connaître les configurations des produits pour chaque groupe de consommateurs.

Modèle de segmentation

Supposons l'existence de T segments pour les membres du jury. Soit $\alpha_{(t)}$ la probabilité qu'un individu quelconque appartienne au groupe t , avec $\alpha_{(t)} \in [0; 1]$ et

$$\sum_{t=1}^T \alpha_{(t)} = 1.$$

On notera $\pi_{ij,t}$ la probabilité que le stimulus i soit préféré au stimulus j pour le segment t . De manière analogue à (1), pour chaque classe t , on peut décomposer cette probabilité suivant le modèle de Bradley :

$$\pi_{ij,t} = \frac{\pi_{i,t}}{\pi_{i,t} + \pi_{j,t}} \quad (4)$$

sous la contrainte : $\pi_{i,t} \in [0; 1]$ et $\sum_{i=1}^n \pi_{i,t} = 1$.

En supposant l'indépendance des observations, la distribution de $Y_h = {}^t(Y_{ij,h}/1 \leq i < j \leq n)$ est un mélange fini de distributions de Bernoulli multivariées de densité :

$$g(y_h; \pi, \alpha) = \sum_{t=1}^T \alpha_{(t)} f(y_h; \pi_{(t)}) \quad (5)$$

où $\pi = {}^t(\pi_{(1)}, \dots, \pi_{(t)})$, $\pi_{(t)} = {}^t(\pi_{1,t}, \dots, \pi_{n,t})$, $\alpha = {}^t(\alpha_{(1)}, \dots, \alpha_{(T)})$ et $f(y_h; \pi_{(t)})$ est définie de façon analogue à (3).

La vraisemblance associée à l'ensemble des observations devient :

$$L(y; \pi, \alpha) = \prod_{h=1}^H g(y_h; \pi, \alpha). \quad (6)$$

Les paramètres π et α sont estimés par maximum de vraisemblance en utilisant un algorithme de type EM (Dempster *et al* 1977).

Estimation des paramètres

Pour mettre en œuvre l'algorithme EM, on introduit la variable aléatoire non observée Z_{ht} , définie par :

$$Z_{ht} = \begin{cases} 0 & \text{si le juge } h \text{ n'appartient pas à la classe } t \\ 1 & \text{si le juge } h \text{ appartient à la classe } t. \end{cases}$$

Soient $Z_h = {}^t(Z_{h1}, \dots, Z_{hT})$ et $Z = {}^t(Z_1, \dots, Z_H)$. On suppose que les variables non observées Z_h sont indépendantes et identiquement distribuées suivant une loi multinomiale de paramètre α dont la densité est définie par :

$$k(Z_h; \alpha) = \prod_{t=1}^T (\alpha_{(t)})^{Z_{ht}}. \quad (7)$$

On suppose de plus que $Y_1, \dots, Y_H$ sachant $Z_1, \dots, Z_H$ sont indépendantes et que la densité de Y_h , conditionnelle à Z_h , est donnée par :

$$q(y_h; \pi/Z_h) = \prod_{t=1}^T [f(y_h; \pi_{(t)})]^{Z_{ht}}. \quad (8)$$

soit si $Z_{ht_0} = 1$, $Z_{ht} = 0$ si $t \neq t_0$, $q(y_h; \pi/Z_h) = f(y_h; \pi_{(t_0)})$.

En considérant les variables Z_h comme des données manquantes, la fonction de vraisemblance associée à l'ensemble des observations est définie à partir de (7) et (8) par (Dempster *et al.*, 1977) :

$$L_c(y, Z; \pi, \alpha) = \prod_{h=1}^H \prod_{t=1}^T [\alpha_{(t)} f(y_h; \pi_{(t)})]^{Z_{ht}} \quad (9)$$

et donc le logarithme de la fonction de vraisemblance définie par (9) s'écrit :

$$\ln L_c(y, Z; \pi, \alpha) = \sum_{h=1}^H \sum_{t=1}^T Z_{ht} \ln \alpha_{(t)} + \sum_{h=1}^H \sum_{t=1}^T Z_{ht} \ln f(y_h; \pi_{(t)}). \quad (10)$$

Etape E : A l'itération s de l'algorithme, les variables Z_h étant non observées, on évalue leur espérance connaissant le vecteur des observations y et les estimations de π et α obtenues à l'itération précédente notées respectivement $\hat{\pi}^{(s-1)}$ et $\hat{\alpha}^{(s-1)}$:

$$E[Z_h/y; \hat{\pi}^{(s-1)}, \hat{\alpha}^{(s-1)}] = \frac{\hat{\alpha}_{(t)}^{(s-1)} [f(y_h; \hat{\pi}_{(t)}^{(s-1)})]}{\sum_{u=1}^T \hat{\alpha}_{(u)}^{(s-1)} [f(y_h; \hat{\pi}_{(u)}^{(s-1)})]} = Z_{ht}^{(s-1)}. \quad (11)$$

Ainsi, pour un individu h , les données non observées z_{ht} sont remplacées par les probabilités *a posteriori* d'appartenance de ce juge aux différents segments.

On évalue donc l'espérance du logarithme de L_c défini par (10) par rapport à Z , de la façon suivante :

$$E_Z[\ln L_c(y; \pi, \alpha/Z)] = \sum_{h=1}^H \sum_{t=1}^T Z_{ht}^{(s-1)} \ln \alpha_{(t)} + \sum_{h=1}^H \sum_{t=1}^T Z_{ht}^{(s-1)} \ln f(y_h; \pi_{(t)}) \quad (12)$$

Etape M : On maximise (12) par rapport à π et α . Les estimations de ces paramètres se font de façon indépendante.

Estimer α revient à maximiser

$$\sum_{h=1}^H \sum_{t=1}^T Z_{ht}^{(s-1)} \ln \alpha_{(t)} \quad (13)$$

sous la contrainte $\sum_{t=1}^T \alpha_{(t)} = 1$.

On obtient alors :

$$\hat{\alpha}_{(t)}^{(s)} = \sum_{h=1}^H \frac{Z_{ht}^{(s-1)}}{H}. \quad (14)$$

Pour estimer π , on maximise $\sum_{h=1}^H \sum_{t=1}^T Z_{ht}^{(s-1)} \ln f(y_h; \pi_{(t)})$. Cette estimation se fera par maximum de vraisemblance en utilisant une méthode itérative proposée par Dysktra (voir présentation en annexe 1).

On réitère les étapes *E* et *M* jusqu'à convergence de l'algorithme. Le critère d'arrêt sera basé sur la différence entre les logarithmes de la vraisemblance définie par (6) évalués en $(\hat{\pi}^{(s)}, \hat{\alpha}^{(s)})$ et $(\hat{\pi}^{(s-1)}, \hat{\alpha}^{(s-1)})$.

Les valeurs initiales $(\hat{\pi}^{(0)}, \hat{\alpha}^{(0)})$ sont choisies aléatoirement.

A l'issue de cette procédure, nous disposons des estimations des configurations des produits, encore appelées scores de Bradley pour chacun des T segments, nous

connaissons les poids des différentes classes, ainsi que les probabilités d'appartenance de chaque individu aux différents segments.

Détermination du nombre de classes homogènes

Nous adopterons une démarche ascendante pour la détermination du nombre de segments homogènes pour les consommateurs. La sélection du nombre de classes sera basée sur la statistique de test du rapport des vraisemblances en utilisant (6) obtenues avec T et $T + 1$ segments et définie par :

$$U_T = -2 \ln \left(\frac{L_{(T)}}{L_{(T+1)}} \right), \quad (15)$$

où $L_{(T)}$ et $L_{(T+1)}$ représentent respectivement les valeurs maximales prises par la vraisemblance pour T et $T + 1$ segments.

L'hypothèse nulle suppose l'existence de T segments de consommateurs.

La statistique définie par (15) ne suivant pas, sous cette hypothèse, une loi de chi deux avec un nombre de degrés de liberté connu (McLachlan et Basford, 1988), une procédure de test de significativité de Monte Carlo est mise en œuvre. Cette méthode nécessite la génération de $K - 1$ échantillons aléatoires de Monte Carlo issus de la configuration sous T segments. Pour chaque série de simulations, le modèle de Bradley pour T et $T + 1$ segments est estimé et la statistique de test U_T est calculée.

On adopte alors la règle de décision suivante : si la valeur de U_T obtenue avec les données observées est supérieure à $K(1 - p^*)$ valeurs de U_T obtenues à l'aide des données simulées, alors l'hypothèse nulle est rejetée au niveau de significativité p^* considéré (Dillon *et al*, 1993).

Application

A titre d'illustration, on traitera des données issues d'une étude d'optimisation de la composition d'un cocktail. Les 7 produits comparés sont fabriqués selon un plan de mélange (Cornell, 1981) en faisant varier les proportions de jus de mangue, de jus de citron et de sirop de cassis (figure 1).

Ces sept mélanges sont présentés à 60 consommateurs lors d'épreuves de comparaisons par paires. Le protocole de présentation des produits, détaillé en annexe 2, est un plan en blocs incomplets.

Les résultats du modèle de segmentation obtenus sur les données de mélange pour $T = 1$ à 4 classes sont présentés dans le tableau 1.

Les simulations réalisées ($K = 100$) permettent de conclure à l'existence de 3 classes latentes dans le jury considéré (tableau 2). La statistique de test observée pour $T = 2$ est significative à 3% alors que pour $T = 3$ on accepte l'hypothèse nulle (figure 2).

La figure 3 donne les configurations des produits dans les différents segments. On retrouve dans la classe 1 des consommateurs préférant des produits faiblement

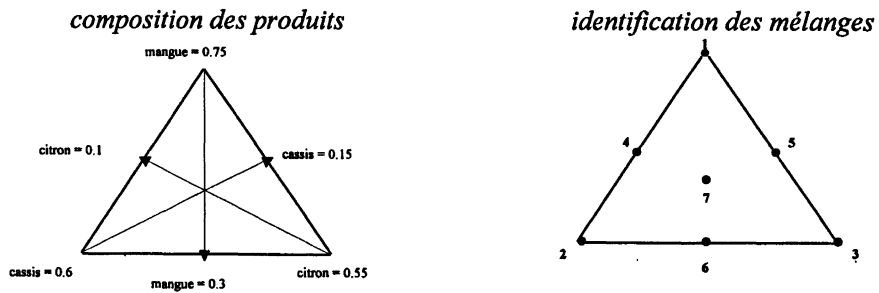


FIGURE 1
Plan de mélange utilisé pour l’élaboration des cocktails

TABLEAU 1
Résultats du modèle de segmentation.
Poids $\alpha(t)$ des segments et configuration $\pi(t)$ des produits

T	t	Poids $\alpha(t)$	Configuration des produits $\pi_i(t)$						
			1	2	3	4	5	6	7
1	1	1.00	0.156	0.144	0.059	0.230	0.095	0.140	0.177
	2	0.65	0.207	0.142	0.005	0.394	0.047	0.068	0.136
2	1	0.35	0.056	0.101	0.272	0.063	0.131	0.234	0.144
	2	0.39	0.092	0.338	0.000	0.477	0.008	0.025	0.061
	3	0.42	0.075	0.119	0.215	0.093	0.111	0.216	0.172
3	1	0.19	0.462	0.000	0.000	0.263	0.125	0.052	0.097
	2	0.19	0.462	0.000	0.000	0.263	0.125	0.052	0.097

TABLEAU 2
Résultats du modèle de segmentation.
Logarithme de la vraisemblance ($\ln L_T$), rapport de vraisemblance (U_T)
et niveau de signification en fonction du nombre T de classes.

	Nombre de classes T			
	1	2	3	4
$\ln L_{(T)}$	-312.18	-286.62	-274.33	-267.56
U_T	51.1	24.6	13.5	
Niveau de signification	0.00	0.03	0.33	

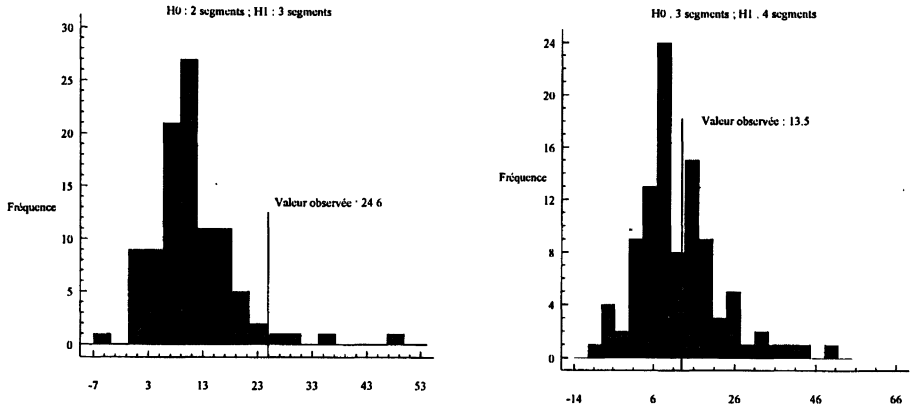


FIGURE 2
Histogrammes des fréquences de répartition de la statistique U_T obtenue à partir des 100 échantillons de Monte Carlo

dosés en citron alors que les membres du segment 3 valorisent la proportion de mangue. La classe 2 discrimine moins bien les produits et semble préférer des mélanges avec peu de jus de mangue.

Les poids des différents segments (respectivement de 0.39, 0.42 et 0.19) peuvent s’interpréter comme les probabilités d’appartenance d’un individu du jury aux trois classes.

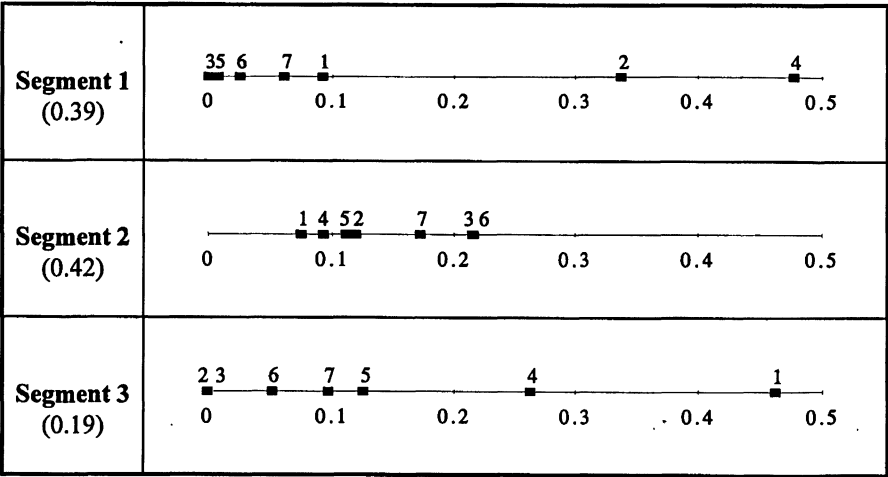


FIGURE 3
Solution pour $T = 3$ classes.
Configuration des produits dans les différents segments et poids de chaque segment.

Conclusion

Les épreuves de comparaisons par paires sont particulièrement adaptées à des mesures hédoniques car elles ne nécessitent pas l'utilisation d'une échelle de notation. La classification du jury de consommateurs (ainsi que le test du nombre de classes latentes) présente un intérêt en marketing et lors d'étude du comportement de consommateurs. La méthode présentée dans ce travail fournit ainsi une aide appréciable pour cerner par exemple l'acceptabilité d'une gamme de produits auprès d'un panel.

L'influence sur les résultats de la présentation des produits en blocs incomplets demande cependant à être précisée. Des études de simulations pourraient permettre de fournir à l'utilisateur des règles de choix de plans expérimentaux pour l'organisation de ce type d'épreuves.

Note : Les programmes utilisés dans ce travail ont été développés en langage SAS/IML (SAS Institute Inc., Cary, NC, USA).

Références

- BRADLEY R.A., TERRY M.E. (1952). Rank analysis of incomplete block designs : The method of paired comparisons. *Biometrika*, 39, 324-45.
- CARROLL J.D. (1980). Models and methods for multidimensional analysis of preferential choice (or other dominance data) in : E.D. Langermann and H. Feger (eds), *Similarity and choice*. Bern : Huber. pp. 234-289.
- COOMBS C.H., (1958). On the use of inconsistency of preferences in psychological measurement. *Journal of experimental psychology* 55, 1-7.
- CORNELL J.A., (1981). Experiments with mixtures. John Wiley & Sons, New York.
- DAVID H.A., (1988). The method of paired comparisons. Oxford University Press, New York.
- DE SARBO W.S., DE SOETE G., ELIASABERG J., (1987). A new stochastic multidimensional unfolding model for the investigation of paired comparison consumer preference/choice data. *Journal of Economical Psychology* 8, 357-384.
- DE SOETE G., CARROLL J.D., (1983). A maximum likelihood method for fitting the wandering vector model. *Psychometrika*, 48, 4, 553-566.
- DE SOETE G., DESARBO W.S., (1991). A latent class probit model for analyzing pick any/N data. *Journal of Classification*, 8, 45-63.
- DEMPSTER A.P., LAIRD N.M., RUBIN D.B., (1977). Maximum likelihood from incomplete data via the EM algorithm. *Journal of the Royal Statistical Society*, B, 39, 1-38.
- DILLON W.R., KUMAR A., SMITH DE BORRERO M., (1993). Capturing individual differences in paired comparisons : an extended BTL model incorporating descriptor variables. *Journal of Marketing Research*, XXX, 42-51.

- DYKSTRA O. JR., (1956). A note on rank analysis of incomplete block designs : A method of paired comparisons employing unequal repetitions on pairs. *Biometrics* 12, 301-306.
- MCLACHLAN G.L., BASFORD K.E., (1988). Mixtures models : inference and application to clustering. Marcel Dekker, New York.
- O'MAHONY M., MASUOKA S., ISHII R., (1994). A theoretical note on difference tests : models, paradoxes and cognitive strategies. *Journal of Sensory Studies* 9, 247-272.
- SCHONEMANN P.H., WANG M.M., (1972). An individual difference model for the multidimensional analysis of preference data. *Psychometrika* 37, 3, 275-309.
- THURSTONE L.L., (1927). A law of comparative judgment. *Psychological Review* 34, 273-286.

Annexes

Annexe 1 : Estimation de $\pi_{(t)}$ (Dykstra, 1956)

Soient :

$$r_{ij,t}^{(s-1)} = \sum_{h=1}^H Z_{ht}^{(s-1)} y_{ij,h} \delta_{ij,h}$$

et

$$R_{ij,t}^{(s-1)} = \sum_{h=1}^H Z_{ht}^{(s-1)} \delta_{ij,h}$$

où $\delta_{ij,h}$ prend la valeur 1 si la paire (i, j) est vue par le juge h et 0 sinon. Estimer $\pi_{(t)}$ revient donc à maximiser :

$$\begin{aligned} w(\pi_{(t)}) &= \sum_{i=1}^{n-1} \sum_{j=i+1}^n r_{ij,t}^{(s-1)} (\ln(\pi_{i,t}) - \ln(\pi_{j,t})) \\ &\quad + \sum_{i=1}^{n-1} \sum_{j=i+1}^n R_{ij,t}^{(s-1)} (\ln(\pi_{j,t}) - \ln(\pi_{i,t} + \pi_{j,t})) \end{aligned}$$

On peut montrer que :

$$\frac{\partial w(\pi_{(t)})}{\partial \pi_{i,t}} = \sum_{j \neq i} \frac{r_{ij,t}^{(s-1)}}{\pi_{i,t}} - \sum_{j \neq i} \frac{R_{ij,t}^{(s-1)}}{(\pi_{i,t} + \pi_{j,t})}$$

Le paramètre $\pi_{i,t}$ est estimé itérativement suivant la méthode proposée par Dykstra (1956); à l'étape u , on a :

$$\hat{\pi}_{i,t}^{(s,u)} = \frac{\sum_{j \neq i} r_{ij,t}^{(s-1)}}{\sum_{j \neq i} \frac{R_{ij,t}^{(s-1)}}{(\hat{\pi}_{i,t}^{(s,u-1)} + \hat{\pi}_{j,t}^{(s,u-1)})}}.$$

On initialise l'algorithme en prenant $\hat{\pi}_{i,t}^{(s,0)} = \hat{\pi}_{i,t}^{(s-1)}$. Le critère d'arrêt est basé sur l'écart entre w évalué en $\hat{\pi}_{i,t}^{(s,u)}$ et $\hat{\pi}_{i,t}^{(s,u-1)}$.

Annexe 2 : Protocole de présentation des produits aux consommateurs

Afin d'éviter des problèmes de saturation, chacun des 60 consommateurs a participé à deux séances de dégustation. Au cours de chacune d'entre elles, un sujet voit 4 produits et évalue 4 paires parmi les 6 possibles. Le produit 7, point central du plan de mélange, est présenté lors des deux sessions et les 3 autres produits sont choisis en utilisant un plan en blocs incomplets équilibrés (10 blocs).

Un consommateur verra par exemple les produits 1, 2, 3, 7 lors de la séance 1 et les produits 4, 5, 6, 7 lors de la séance 2.

Lors d'une session, la présentation par paires des 4 produits est faite selon un plan de type «linked paired-comparison design», dont la construction est basée sur des plans en blocs incomplets équilibrés (David, 1988). A titre d'exemple, les 6 modes de présentation des produits 1, 2, 3, 7 seront :

(1,2); (2,7); (7,3); (3,1)
 (2,3); (3,1); (1,7); (7,2)
 (3,7); (7,2); (2,1); (1,3)
 (7,1); (1,3); (3,2); (2,7)
 (1,2); (2,3); (3,7); (7,1)
 (7,3); (3,2); (2,1); (1,7)

Ce type de plan équilibre le nombre d'apparitions de chacune des paires ainsi que les effets d'ordre de présentation des produits dans chaque paire.

Pour l'ensemble du protocole de dégustation, les produits (hormis le produit 7) sont donc vus 120 fois. Le produit 7, présent lors des deux séances, est vu 240 fois. Les 15 paires ne comportant pas le produit 7 sont évaluées 16 fois et les 6 paires comportant le produit 7 le sont 40 fois.

Nous disposons donc de 480 comparaisons par paires sur les 1260 possibles dans le cas d'un plan complet (les 21 paires vues par les 60 consommateurs).