

01 MAG
1999-05-31
CA
2627

Essays in International Market Segmentation

Frenkel ter Hofstede

Stellingen

1. Het negeren van de impliciet gestratificeerde steekproef designs voor internationale marktsegmentatie kan leiden tot verkeerde inzichten in de structuur van internationale marktsegmenten en derhalve tot onjuiste beslissingen van managers.
dit proefschrift
2. In tegenstelling tot laddering, de traditionele methode voor het meten van means-end chain relaties van consumenten, is de associatie patroon techniek zeer geschikt voor gebruik in representatieve internationale steekproeven.
dit proefschrift
3. Internationale marktsegmentatie op basis van means-end chain relaties van consumenten is een belangrijk middel voor het formuleren van internationale geïntegreerde segmentatie strategieën voor productontwikkeling en communicatie.
dit proefschrift
4. Bayesiaanse modelformuleringen en Markov Chain Monte Carlo schattingstechnieken kunnen dienen ter verbetering van de geografische configuratie van internationale marktsegmenten, hetgeen de toegankelijkheid van deze segmenten verhoogt, zonder dat daarbij de reactiviteit van de segmenten ernstig wordt geschaad.
dit proefschrift
5. Resultaten van internationale studies verkregen middels psychometrische meetinstrumenten kunnen ernstig verstoord worden door verschillen in schaalgebruik van consumenten tussen landen, alsmede door verschillen in schaalgebruik tussen consumenten binnen eenzelfde land.
dit proefschrift
6. De accumulatie van kennis over het bestaan en de aard van transnationale en "global" segmenten wordt in belangrijke mate gehinderd door een gebrek aan aandacht voor de methodologische aspecten van internationale marktsegmentatie bij marktonderzoekers en marketeers.
7. Tijdsduur aggregatie kan de nauwkeurigheid van parameterschattingen van parametrische proportionele hazard modellen ernstig schaden, ongeacht of de modellen geformuleerd worden in continue dan wel discrete tijd.

F. ter Hofstede en M. Wedel, 1998, A Monte Carlo study of time-aggregation in continuous-time and discrete-time parametric hazard models, Economics Letters, 58 (2), 149-156.

8. De multicollineariteit van parameters in Box-Cox specificaties van baseline hazard functies en het gebruik van de gemiddelde kwadratische fout als criterium voor parameter nauwkeurigheid, leiden tot een overschatting van het effect van tijdsduur aggregatie op de onnauwkeurigheid van de baseline hazard.

F. ter Hofstede en M. Wedel, 1999, Time aggregation effects on the baseline of continuous-time and discrete-time hazard models, Economics Letters, in druk.

9. De nauwkeurigheid van voorspellingen van merkkeuze en tijdstip van productaankopen kan aanzienlijk verbeterd worden door rekening te houden met zowel proportionele als niet-proportionele effecten van prijspromoties, alsmede de heterogeniteit van deze effecten tussen consumenten.

M. Wedel, W.A. Kamakura, W.S. DeSarbo en F. ter Hofstede, 1995. Implications for asymmetry, nonproportionality, and heterogeneity in brand switching from piece-wise exponential mixture hazard models, Journal of Marketing Research, 32 (november), 457-462.

10. Als gevolg van sociale versterkingsmechanismen kunnen culturele variabelen een modererende rol spelen in de relatie tussen centrale disposities en de innovativiteit van consumenten.

J.E.B.M Steenkamp, F. ter Hofstede en M. Wedel, 1999, A cross-national investigation into the individual and national-cultural antecedents of consumer innovativeness," Journal of Marketing, 63 (april), in druk.

11. Mens en computer zijn in hoge mate complementair aangezien computers uitblinken in het uitwerken van specifieke gevallen aan de hand van algemene principes, terwijl de mens een sterk ontwikkeld vermogen heeft om algemene principes af te leiden uit specifieke gevallen.

12. Het Spaanse verkeersbord met de vermelding "neemt u de verkeersborden serieus," duidt op realiteitszin, maar niet op logisch deductieve vaardigheden van haar bedenkers.

13. Het feit dat een groot deel van de Nederlandse promovendi exact twaalf stellingen formuleert zonder dat promotiereglementen daartoe aanleiding geven, kan duiden op zowel bijgelovigheid als behoudendheid van aanstaande wetenschappers en is derhalve een zorgwekkende indicatie van de verdere ontwikkeling van de wetenschap in Nederland.

Stellingen behorende bij het proefschrift
Essays in International Market Segmentation
Frenkel ter Hofstede
9 juni 1999

Essays in International Market Segmentation



Promotoren: Prof. dr ir J.E.B.M. Steenkamp
bijzonder hoogleraar methoden en technieken van
marktonderzoek met betrekking tot de Europese
consument van agrarische producten aan de
Landbouwniversiteit Wageningen en gewoon
hoogleraar marketing aan de Katholieke Universiteit
van Leuven

Prof. dr M. Wedel
hoogleraar marktonderzoek aan de Rijksuniversiteit
Groningen

Essays in International Market Segmentation

Frenkel ter Hofstede

Proefschrift

ter verkrijging van de graad van doctor
op gezag van de rector magnificus
van de Landbouwwuniversiteit Wageningen
dr. C.M. Karssen,
in het openbaar te verdedigen
op woensdag 9 juni 1999
des namiddags te vier uur in de Aula.

nn 964117

ISBN: 90-5808-064-1

BIBLIOTHEEK
LANDBOUWUNIVERSITEIT
WAGENINGEN

To Judith, Emielou, and Blanche

Contents

ACKNOWLEDGEMENTS	VII
CHAPTER 1 INTRODUCTION	1
1.1 Introduction	1
1.2 Current Status of International Market Segmentation	3
1.2.1 International Geographic Segmentation Studies	7
1.2.2 International Consumer Segmentation Studies	10
1.3 Unresolved Issues	13
1.4 Outline and Objectives	15
CHAPTER 2 SAMPLING DESIGNS IN INTERNATIONAL SEGMENTATION	19
2.1 Introduction	19
2.2 Sample Design and the Mixture Approach	21
2.2.1 Pseudo-Maximum Likelihood Approach	21
2.2.2 Stratified Sampling	24
2.2.3 Cluster Sampling	25
2.2.4 Two-Stage Sampling	25
2.3 Statistical Inference	26
2.3.1 Estimation	26
2.3.2 Asymptotic Standard Errors of the Estimates	26
2.3.3 Criteria for Selecting the Number of Segments	27

2.4 Illustration	28
2.4.1 Background	28
2.4.2 Data	30
2.4.3 The International Value Segmentation Model	31
2.4.4 Comparison of ML and PML Results	32
2.4.5 Substantive Interpretation of the Five-segment PML Solution	37
2.5 Conclusions	38

CHAPTER 3 AN INVESTIGATION INTO THE ASSOCIATION PATTERN TECHNIQUE

41

3.1 Introduction	41
3.2 Techniques for Measuring Means-End Chains	43
3.2.1 Laddering	43
3.2.2 The Association Pattern Technique	45
3.3 Key Research Issues	48
3.3.1 Assumption of Conditional Independence.	48
3.3.2 Between-method Convergent Validity	49
3.4 Method	52
3.4.1 Data	52
3.4.2 Analysis	53
3.5 Results	58
3.5.1 Conditional Independence	58
3.5.2 Convergent Validity	60
3.6 Discussion and Conclusions	63

CHAPTER 4 INTERNATIONAL MARKET SEGMENTATION BASED ON CONSUMER-PRODUCT RELATIONS

67

4.1 Introduction	67
4.2 The International Segmentation Basis	70
4.3 The International Segmentation Methodology	72
4.3.1 Model Formulation	72
4.3.2 Estimation and Model Comparison	76
4.3.3 Country-specific Entropy	77
4.4 Monte Carlo Analysis	78
4.5 Empirical Application	82
4.5.1 Data Collection	82

4.5.2 Results	85
4.5.3 Segment Profiles	93
4.5.4 Predictive Validity of the Model	95
4.5.5 Comparison with Other Methods	95
4.6 Discussion	97
Appendix 4.A	101
Appendix 4.B	105

**CHAPTER 5 A BAYESIAN APPROACH TO IDENTIFYING GEOGRAPHIC
TARGET MARKETS** 107

5.1 Introduction	107
5.2 Model Development and Estimation	110
5.2.1 Previous Approaches to Constrained Segmentation	110
5.2.2 A Response-based Geographic Segmentation Model	112
5.2.3 Model Estimation and Determination of Number of Segments	117
5.3 Synthetic Data Analysis	119
5.4 Empirical Application	122
5.4.1 Background of the Study	122
5.4.2 Data	123
5.4.3 Results	125
5.4.4 Segment Profiles	133
5.4.5 Responsiveness of the Geographic Segments	135
5.5 Discussion	137

CHAPTER 6 DISCUSSION 141

6.1 Summary and Conclusions	141
6.1.1 Two Directions in International Segmentation Research	141
6.1.2 Methodological Issues in International Market Segmentation	144
6.2 Limitations and Issues for Future Research	146

REFERENCES 149

INDEX 165

SAMENVATTING	173
--------------	-----

CURRICULUM VITAE	179
------------------	-----

Acknowledgments

I am indebted to the many people who have contributed to this dissertation project. First and foremost, I would like to express my deep gratitude to my dissertation advisors, Prof. dr ir Jan-Benedict E.M. Steenkamp and Prof. dr Michel Wedel. I admire their knowledge, wisdom, creativity, and research productivity and am forever indebted to them for their continual and extensive guidance and support. Not only did they provide valuable suggestions and comments on parts of the manuscript in various stages of its development, they also contributed directly to this work through their co-authorship of the papers that cover part of this thesis. Our cooperating has been a special privilege. I would like to thank Michel Wedel for introducing me as an undergraduate student to the world of marketing academia, which generated my initial interest in scientific research. I would like to thank Jan-Benedict Steenkamp for the considerable effort he has put into my academic development. Each time he visited Wageningen was a stimulating and valuable experience.

I owe many thanks to Ms. Anke Audenaert for our pleasant and stimulating cooperation. She has put much effort in collecting the data that were used in chapter 3 and is a co-author of the paper on which that chapter is based. Many thanks to Prof. dr ir Marnik G. Dekimpe, Prof. dr ir Matthieu T.G. Meulenberg, and Prof. dr ir Hans C.M. Van Trijp for their valuable comments. I have received much help and encouragement

from the discussions with fellow Ph.D. students and researchers, particularly dr Youngchan Kim, dr ir Joost M.E. Pennings, and ir Peeter W.J. Verlegh.

This thesis benefited from the many comments and suggestions made by editors and reviewers of the papers that are presented in chapters 2, 3, and 4. This thesis also benefited from feedback during presentations at Carnegie Mellon University, Catholic University of Leuven, New York University, Pennsylvania State University, the University of California at Los Angeles, the University of Colorado at Boulder, the University of Groningen, the University of Iowa, the University of Michigan, the University of Pennsylvania, and Stanford University. Special thanks to the members of the W.T.G., dr Tammo H.A. Bijmolt, dr ir Joost M.E. Pennings, and dr Hiek R. van der Scheer, for the stimulating discussions, which, although indirectly, undeniably enhanced this thesis.

The Commission of the European Communities financially supported the surveys that were used in this thesis, grant no. AIR2-CT94-1066. Additional financial aid of the University of Groningen and Wageningen University is gratefully acknowledged. I am grateful to Mr. Jaap A. Bijkerk for his valuable work in preparing the data for analysis and creating the graphical exhibits presented in this thesis. Many thanks to Ms. W. Christa Biervliet for editorial support and to Mr. Eric Gersons for research assistance. I would like to thank ir Koert van Ittersum for coding the geographic data for the analysis in chapter 5 and Ms. Ellen C. Vossen and Mrs. Liesbeth M. Hijwegen for their secretarial support.

Finally, I would like to extend my deep appreciation to my wonderful wife Judith for her continuing confidence and support. Not only has she contributed to this work by word processing parts of this thesis, she has also supported me mentally during its compilation. I dedicate this thesis to her and our daughters Emielou and Blanche.

Wageningen, March 1999

Chapter 1

Introduction

1.1 Introduction

International market segmentation has become an important issue in developing, positioning, and selling products across national borders. International segmentation helps companies obtain an appropriate positioning of their products across borders and to target potential customers at the international segment-level. A basic challenge for companies is to effectively deal with the structure of heterogeneity in consumer needs and wants across borders and to target groups or segments of consumers in different countries. These segments reflect geographic groupings or groups of individuals and consist of potential consumers who are likely to exhibit similar responses to marketing efforts.

A traditional and natural form of international segmentation is to adopt a multi-domestic strategy where each country represents a separate segment. A multi-domestic strategy amounts to the selection of countries on the basis of their local advantages. In such an approach no coordination is required between countries, products are produced locally and are tailored to satisfy local needs. Distinct advertising, distribution,

and pricing strategies are developed for targeting consumers in each country, and competition is managed at a national level. Competitive moves to fight competitors are performed on a country-to-country basis and do not take the developments in the rest of the world into account. Thus, in segmenting their markets, firms operating according to a multi-domestic approach can suffice with the standard segmentation techniques that are developed for domestic markets. International segmentation becomes a particularly relevant issue when companies adopt a global or pan-regional strategy, that is, a strategy integrated across national borders.

In many industries, national borders are becoming less and less important to delineate international activities, rendering multi-domestic strategies less effective. Developments accelerating this trend include rapidly falling national boundaries, regional unification, standardization of manufacturing techniques, global investment and production strategies, expansion of world travel, rapid increase in education and literacy levels, growing urbanization among developing countries, free flow of information, labor, money, and technology across borders, increased consumer sophistication and purchasing power, advances in telecommunication technologies, and the emergence of global media (Hassan and Katsanis 1994; Alden, Steenkamp, and Batra 1999; Mahajan and Muller 1994). Many global companies such as Coca-Cola, McDonalds, Sony, British Airways, Ikea, Toyota, and Levi-Strauss have successfully integrated their international strategies. By globalizing their strategies such companies benefit from several advantages, including cost reductions through economies of scale, improved quality of products, and increased competitive power (Levitt 1983; Yip 1995).

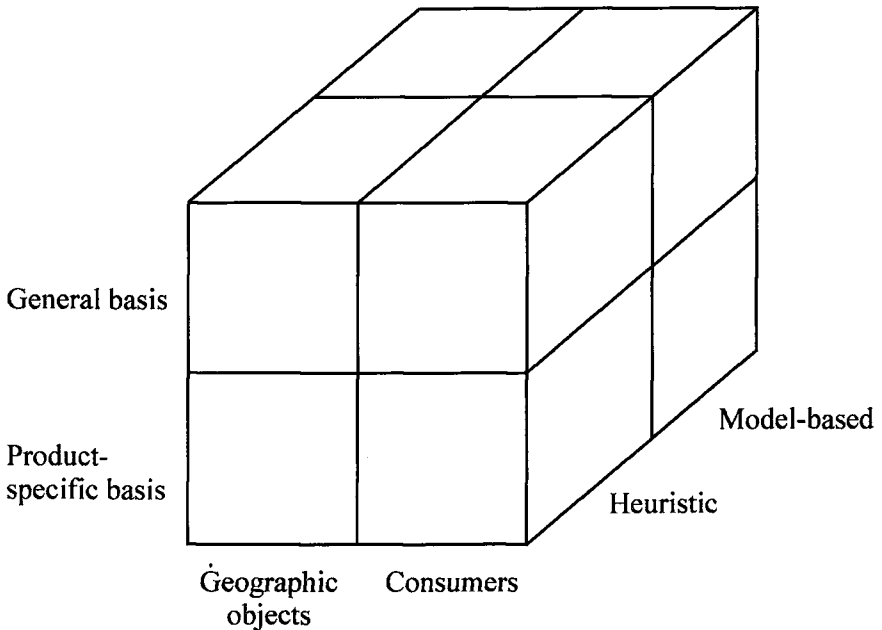
Still, companies cannot serve the entire world population with fully standardized marketing strategies. A major challenge facing international marketers is to identify global market segments and reach these segments with tailored products and marketing programs (Hassan and Katsanis 1994). Many companies recognize that groups of consumers in different countries often have more in common with one another than with other consumers in the same country. Hence, they choose to serve segments that transcend national borders. For example, Perrier and BMW are

successfully targeting their products to an international segment of cosmopolitan consumers and Sony and Levy Strauss developed particular products and marketing programs for teenagers worldwide (Hassan and Katsanis 1994).

International segmentation aids in structuring the heterogeneity that exists among consumers and nations and helps to identify segments that can be targeted in an effective and efficient way. Several studies have looked into the concept of international market segmentation. This chapter provides an overview of the current status of international market segmentation research and segmentation methodologies. Previous literature in the area is discussed by means of a framework for classifying international segmentation. Finally, an outline and objectives of the thesis are provided.

1.2 Current Status of International Market Segmentation

International segmentation can be classified along three dimensions that affect the effectiveness of international marketing strategies (see Figure 1.1). The first dimension is related to the basis for international segmentation, which is *general* or *product-specific* (Wedel and Kamakura 1998). The segmentation basis is a set of characteristics that defines the segments. General bases are independent of product characteristics, whereas product-specific bases depend on the particular product under study. Examples of general international segmentation bases are geographic locations, such as regions or countries, cultural characteristics, economic indicators, political characteristics, demographics, consumer values, and life-styles. Product-specific bases typically rely on import statistics, brand penetration rates, economic and legal constraints, importances of attributes and benefits, purchase intentions, and preferences.

*Figure 1.1**Three dimensions of international segmentation*

Segments identified with product-specific bases such as attribute importances tend to be more effective when responsive segments are required, that is, segments that respond uniquely to marketing efforts (cf., Wedel and Kamakura 1998). Responsiveness of international segments is closely related to the issue of globalization and of standardization versus adaptation of elements of the marketing mix (Farley and Lehman 1994). Initiated by Levitt (1983), the issue of whether to standardize products and marketing mixes internationally versus adapting them to local tastes has been a central topic of debate in the international marketing literature (Szymanski, Bharadwaj, and Varadarajan 1993). Responsive international or global segments provide a way to target consumers with standardized products and marketing mixes across national borders.

Among the set of general segmentation bases, values are of particular interest for international segmentation. Values are independent of the product but closer to the consumer and are particularly actionable

for segmenting the market when behavioral patterns across one or more product categories are the focal point of interest (Kamakura et al. 1993). From a managerial perspective, patterns of behavior are of special interest when the firm offers multiple product categories (e.g., soaps, cosmetics, etc.), which relate to broader behavioral domains (e.g., personal care).

The second dimension in Figure 1.1 is related to the level of aggregation. International segments may be defined in terms of *geographic objects* or in terms of *consumers*,¹ which result in geographic segments and consumer segments, respectively. Geographic segments typically cover a single country, groups of countries, regions within a number of countries, or larger geographic areas such as the Benelux, the European Union (EU), or the North Atlantic Free Trade Area (NAFTA). Geographic segmentation plays an important role when firms pursue a geographic differentiation strategy. For example, Philips basically treated the geographic segments of northern and southern Europe differently when marketing its personal care products and Unilever targets different geographic regions with different products and brands. As opposed to geography, international segments may be defined in terms of consumers. Such segments may consist of consumers of different nationalities that share similar needs and provide opportunities for targeting a limited set of segments that exist in many countries.

Geographic segments could prove to be effective due to their geographic locations. A rationale for entering geographic areas is that it results in accessible and costs effective strategies through centralization of activities such as production, sales force management, service support, logistics, advertising, production of promotional materials, and training of personnel (Tackeuchi and Porter 1986; Yip 1995). On the other hand, geographic segments may overlook the differences that exist between consumers in these countries, affecting the responsiveness of segments. International consumer segmentation may provide a better way to identify segments that are more similar in terms of their needs. Still, targeting a consumer segment that exists in many countries may not

¹ Whereas the discussion concentrates on consumers, the formulation generalizes to decision-makers, including consumers and industrial customers.

always be very cost efficient from a logistics perspective. Especially in industries where distribution costs constitute a large part of the total costs, such as in retailing and in industries dealing with perishable products, geographically dispersed consumer segments will often not allow profitable entry strategies to be pursued.

The third dimension is related to the method used to identify international segments, which is either *heuristic* or *model-based*. Heuristic approaches to international segmentation rely on heuristics such as subjective choice rules to select geographic target markets or cluster and Q-factor analysis tools. These approaches are not based on theory and tend to provide less guidance for formulating marketing strategies. Q-factor analysis has been demonstrated not to be an appropriate technique for segmentation (Dillon and Goldstein 1984; Stewart 1981). Cluster analysis is more appropriate, but also has a number of limitations. It generates deterministic classifications often based on subjective optimization criteria. Cluster analysis does not fit within the framework of standard statistical theory and does not provide reliability judgements of the results.

The current state of the art methodology in segmentation shows an increased use of model-based approaches to segmentation. Model-based segmentation identifies segments based on a particular representation of reality. A model-based approach allows one to identify segments using theories of consumer behavior and to incorporate particular managerial considerations that are relevant for international firms. *Finite mixture models* are model-based approaches to segmentation and overcome the limitations of cluster analysis. They allow hypothesis testing and estimation according to basic statistical principles. Finite mixture models simultaneously identify segments and unknown parameters that characterize such segments. Wedel and DeSarbo (1995) have proposed a general class of mixture regression models, based on generalized linear models. The models allow the identification of response-based segments, where a dependent variable is distributed according to one of the members of the exponential family, including the binomial, Poisson, negative-binomial, normal, gamma, exponential and Dirichlet distributions. A limitation of the mixture approach to segmentation is that

it is not easily implemented and may suffer from local optima in estimation.

Finite mixture models extend to the *Bayesian framework* (Diebolt and Robert 1994; Hoijsink and Molenaar 1997). The current upsurge in Bayesian statistics ensuing from the breakthrough in computational methods has inspired researchers to solve complex problems in different fields of research, including marketing (e.g., Allenby and Ginter 1995; Lenk et al. 1996; Rossi and Allenby 1993). In a Bayesian framework, parameters are not fixed but are random quantities and the objective is to estimate the posterior distributions of unknown parameters, given the data. Bayesian mixture models are estimated using Markov Chain Monte Carlo (MCMC; cf., Gelman et al. 1995) methods, which involve integration over the posterior distribution of the parameters by iteratively drawing samples from the full conditional distributions of the model parameters, given the data. Sample statistics such as the mean, mode, and other percentiles are then computed from the draws to characterize the posterior distribution. Such a sampling-based approach is very flexible, it is convenient for imposing particular restrictions on the model, and easily allows incorporating within-segment heterogeneity and prediction of parameters at the individual level (cf., Allenby, Arora, and Ginter 1998; Allenby and Rossi 1999; Lenk and DeSarbo 1999). The flexibility of MCMC methods allows estimation of complex international segmentation models. Still, such methods are not easily implemented and may suffer from trapping states that may occur in MCMC estimation.

The following sections will provide an overview of previous studies in international segmentation and pay particular attention to the three dimensions depicted in Figure 1.1. Studies that identify geographic international segments will be discussed first, followed by consumer segmentation studies.

1.2.1 International Geographic Segmentation Studies

Quite a few studies have identified geographic segments based on general macro-level data. For example, Sethi (1971) demonstrated the feasibility of collapsing a large set of general, multi-country data into country segments. A heuristic approach is taken, using a general segmentation

basis of 29 political, socioeconomic, trade, transportation, communication, biological, and personal consumption variables obtained for 91 countries. He identified four factors: aggregate production and transportation, personal consumption, trade, health, and education. The countries were cluster analyzed based on their factor scores, resulting in seven country segments. According to Sethi, the development of a successful international strategy depends to a large extent on a firms' ability to segment its world markets such that uniform sets of marketing decisions can be applied to a group of countries.

Day, Fox, and Huszagh (1988) partly replicated the study of Sethi (1971) with the primary purpose to identify global industrial market segments. For 96 countries, 18 indicators of economic development (viz. demographic, economic, health, educational, and transportation variables) were factor and cluster analyzed similar to the analysis of Sethi (1971). Six segments were revealed of which only one corresponded to a segment of Sethi's study. The disparity of the segment structure does not confirm the stability of such general segments across time, but may also result from the instability of the heuristic methods used to identify the segments.

Another approach is to classify countries on general variables that are related to culture, mostly concerning employee work attitudes (e.g., Sirota and Greenwood 1971; Ronen and Kraut 1977). Widely known is the seminal work of Hofstede (1980), who derived four main conceptual dimensions on which cultures exhibit significant differences, viz. individualism/collectivism, power distance, masculinity/femininity, and uncertainty avoidance. Hofstede clustered countries on the basis of their placement on the four cultural dimensions which revealed eight country segments. Such segments are relevant for several purposes, including the design of sales force stimulation systems across different cultures (Usunier 1996). In a review of eight studies, Ronen and Shenkar (1985) concluded that the area of country segmentation based on culture data is lacking methodological rigor. An interesting observation is that in all eight studies the countries tend to group together in the same geographic segments. However, more recently Kahle (1995) performed a cluster

analysis on the Hofstede ratings of 17 countries in Europe and found only one segment replicating Hofstede's clusters.

Vandermerwe and L'Huillier (1989) clustered 173 regions of 18 European countries on five geographic and sociodemographic variables, i.e., language, latitude, longitude, age groups, and income. Not surprisingly, given the strong geographic emphasis of the segmentation basis, they identified geographically contiguous segments. These segments, however, transcended national borders, which supports the existence of segments of Euro-consumers that transcend national borders.

Basically, groupings of geographic objects identified with economic indicators allow marketers to assess what types of products and technologies could be sold in which locations (Day, Fox, and Huszagh 1988). The main focus has been on grouping countries. Vandermerwe and L'Huillier (1989) increased the precision of geographic segments by using smaller regions within countries, but the sociodemographic basis of segmentation will not always lead to very actionable marketing strategies.

Since general bases do not provide much information on product-specific behavior, several studies have identified international segments that are directly related to products or product classes. Askegaard and Madsen (1998) explored the cultural boundaries in Europe by clustering 79 regions from 15 European countries into larger areas. They applied hierarchical cluster analysis to data collected in a food-related lifestyle survey. Twelve contiguous segments were identified that followed clear language borders (e.g., the Dutch speaking part of Belgium was united with the Netherlands). As indicated by the authors, "the regularity of the language/clustering pattern may indicate a fundamental, systematic bias in the translation of the questionnaires." Another problem associated with the application of lifestyle instruments for international segmentation is that lifestyles are rather situation specific. Similar lifestyle types in different cultures may lead to different expressions of behavior (Douglas and Urban 1977; Eshgi and Sheth 1985; Steenkamp 1992).

Helsen, Jedidi, and DeSarbo (1993) assessed the merits of general country segmentation schemes to gain insights into multinational diffusion patterns. The authors developed a model-based mixture methodology that identifies country segments on the basis of diffusion

parameters of the Bass (1969) model and at the same time estimates these parameters within segments. The model is applied to sales data of three products (video recorders, televisions, and CD players) in 12 countries. They compared diffusion-based country segmentation to general country segmentation similar to Sethi (1971). This comparison demonstrated a lack of congruence between the two segmentation schemes. In addition, the results showed strong differences among the products' diffusion-based segmentations, both in terms of number of segments and their composition. This suggests that general country segments may provide little guidance for successful international product introduction.

1.2.2 International Consumer Segmentation Studies

Boote (1983) proposed lifestyles as a basis for identifying international consumer segments for developing cross-national advertising strategies. He applied Q-factor analysis to life-style importance ratings of about 500 women in France, the United Kingdom, and Germany. Four cross-national segments were identified and labeled as 'traditional homemaker,' 'spontaneous,' 'contemporary homemaker,' and 'appearance conscious'. The 'traditional homemaker' and 'contemporary homemaker' segments covered all countries in substantial proportions. The other two segments were dominated by France and Germany, respectively. Cross-national lifestyle segmentation schemes have been developed in the industry as well. For example, Goodyear Tire tailored different marketing mixes and promotional programs to six cross-national lifestyle segments and Ogilvy and Mather identified lifestyle typologies across countries (Hassan and Katsanis 1994).

Kamakura et al. (1993) looked into the issue of formulating standardized pan-European strategies. Building upon the work of Kamakura and Mazzon (1991), the authors developed a rank ordered mixture model for cross-national value segmentation research. Values are very suitable for international segmentation, since they are universal across nations, stable and centrally held by consumers (cf., Schwartz 1992). The sample consisted of about 1500 female consumers in three major countries of the European Union, the United Kingdom, Italy, and Germany. Using rankings of both instrumental and terminal values from

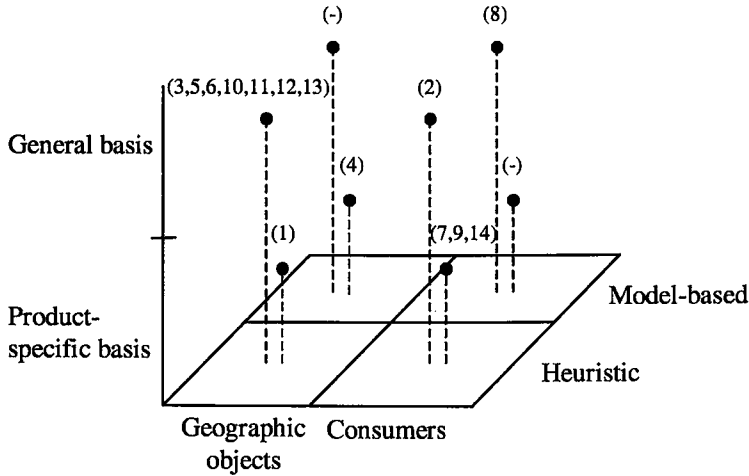
the Rokeach (1973) value list, the authors assessed whether cross national segments exist, which consumers share similar value structures. Five segments were found with different value systems, and were labeled as the 'mature,' 'security,' 'enjoyment,' 'self-directed,' and 'conformity' segments. Although the segments demonstrated strong country effects, still a large part of the consumers did not belong to a segment dominated by one country.

Kale and Sudharsan (1987) recognized the lack of responsiveness of country segmentation schemes, which ignore the within-country heterogeneity and encourage misleading national stereotyping. They proposed a sequential procedure for "strategically equivalent segmentation," which reveals cross-national segments of consumers that are likely to respond similarly to a marketing mix. They suggest to start narrowing down countries to a number of viable entry candidates, based on macro-level data. In a second step, they propose to segment consumers in each country separately and seek for similarities among segments across countries using judgmental evaluations. A thorough theoretical motivation of the procedure, however, is lacking, as well as an empirical application and validation. Moreover, the approach is piecemeal and rather heuristic. Frank, Massy, and Wind (1972) and Wind and Douglas (1972) proposed a similar approach, with the distinction that in the second stage consumer segments should be constructed for country segments instead of countries.

Yavas, Verhage, and Green (1992) viewed transnational consumer segmentation as a "middleground" strategy between absolute standardization across countries and strategies that are adapted to each country separately. They assessed the segment structure of consumers from 6 countries in different continents. Using K-means clustering, four cross-national segments were identified on the basis of perceived risk and brand loyalty for two products, bath soap and toothpaste. The variation across segments appeared to be larger than the variation across countries, which suggests that it is possible to tailor a standardized strategy to consumers in the same international segment in different countries. A major concern is that the samples used for the analyses were highly selective (exclusively middle income women residing in major urban

centers), which affects the representativeness of the findings and may lead to an underestimation of the amount of heterogeneity in these countries. After all, the sample on itself already defines a particular segment.

In a similar vein, Moskowitz and Rabino (1994) bring up the concept of “sensory segmentation” to tailor products to a limited number of cross-national consumer segments as opposed to tailoring products to different countries. If a limited number of consumer segments can be targeted with products satisfying similar needs, global companies serving many countries can reduce the number of products. Based on six ingredients, 29 product prototypes were constructed and tested in a panel of about 800 consumers in four countries from different parts of the world. The products were rated on the basis of a number of sensory product attributes and purchase intentions. By subsequently applying individual-level regressions, factor analysis, and cluster analysis the authors arrived at three segments that span national borders. For each of these segments the methodology predicts optimal products in terms of their attractiveness. The segments are responsive with respect to product characteristics and directly guide product development. The procedure neglects, however, other elements of the marketing mix (such as advertising) since it does not provide any information on specific characteristics of consumers belonging to the same segment. Another serious concern is the limited number of degrees of freedom for the regressions and the piecemeal and heuristic approach to the segmentation problem.

*Figure 1.2**Positions of studies in the international segmentation classification*

- | | |
|---------------------------------------|---------------------------------------|
| 1. Askegaard and Madsen (1998) | 8. Kamakura et al. (1993) |
| 2. Boote (1983) | 9. Moskowitz and Rabino (1994) |
| 3. Day, Fox, and Huszagh (1988) | 10. Ronen and Kraut (1977) |
| 4. Helsen, Jedidi, and DeSarbo (1993) | 11. Sethi (1971) |
| 5. Hofstede (1980) | 12. Sirota and Greenwood (1971) |
| 6. Kale (1995) | 13. Vandermerwe and L'Huillier (1989) |
| 7. Kale and Sudharsan (1987) | 14. Yavas, Green and Verhage (1992) |

1.3 Unresolved Issues

The international segmentation studies previously described are depicted in Figure 1.2. The figure displays a concentration of geographic segmentation studies that use general bases. Model-based approaches to international segmentation are presented in the studies of Kamakura et al. (1993) and Helsen, Jedidi, and DeSarbo (1993), who used mixture model formulations. However, most international segmentation studies rely on heuristic approaches such as cluster or Q-factor analysis tools. In addition, in the case of diffusion patterns across countries, Helsen, Jedidi, and DeSarbo (1993) demonstrated that heuristic approaches are not always appropriate to approximate model-based approaches. As noted by

Jain (1989) there is a need for rigorous techniques for conducting international market segmentation. New developments in model-based segmentation methods may contribute to more accurate and broader applications of international segmentation.

An important issue in international marketing research is that measurement instruments are subject to *response tendencies*. Response tendencies may interfere with the comparability of instruments across cultures because those tendencies may differ across countries (Baumgartner and Steenkamp 1999; Douglas and Craig 1999; Steenkamp, ter Hofstede, and Wedel 1999). Cross-national response bias may hamper the identification of international consumer segments and may lead to incorrect inferences concerning the segmentation structure identified. The international segmentation studies discussed above did not accommodate cross-national response bias, which may be one of the reasons for the identification of country-level or language related segments. In segmentation studies seeking consumer segments that cross national boundaries, response bias may favor the identification of national segments. Indirectly, through the differences in response behavior, nations become an implicit part of the segmentation basis. Clearly, there is a need for international segmentation methods that correct for cross-national response bias. Such models should allow the assessment of how, if, and to what extent response behavior influences the identification of international segments.

Related to differences in response tendencies is the *comparability of primary and secondary data*, which is important for making correct inferences across international borders. In particular in cases where primary consumer data are collected, international segmentation studies do not take much care in reporting on the particular actions undertaken to warrant the quality of the data. To collect international survey data, for example, back translation procedures are required to ensure a similar content of statements across different languages (Brislin 1970). International segmentation studies tend to use selective samples, which reduces heterogeneity in samples, interferes with the representativeness of the results, and may reinforce country effects in segmentation (Kamakura et al 1993; Askegaard and Madsen 1998). The issue of cross-

national comparability is also important for geographic segmentation, which typically relies on secondary data. Secondary data have become widely available but tend to suffer from differences in definitions, incompleteness, and lack of availability across time and space (Ronen and Shenkar 1985). It is of significant importance that international segmentation studies pay more attention to the comparability of measurement instruments in cross-national data collection. In addition, to obtain valid and reliable results, international segmentation should be based on validated measurement instruments that allow standardized procedures in collecting large and representative samples across countries.

Another serious concern is the *sampling design* of international segmentation studies. The data collection procedure typically used in international marketing research leads to national sample sizes that are not proportional to population sizes. Such samples are stratified and not representative for the international population under study. Therefore, international segmentation studies should take this implicit stratification procedure into account. The studies described in section 1.2 do not accommodate sampling designs, nor do they recognize this problem. In addition, some of the studies rely on very complicated sampling designs, such as mixed two stage sampling designs, with stratified sampling across countries and within each country a different sampling procedure. For example, Yavas, Verhage, and Green (1992) used a mixed design of stratified, random, and snowball sampling procedures. Since the segmentation literature does not provide a direct answer to this problem, it is imperative to assess the effects of not accommodating international sampling designs in segmentation and to provide a methodology for accommodating such designs.

1.4 Outline and Objectives

The objective of this thesis is to propose, develop and validate new methodologies for international segmentation to improve the effectiveness of international marketing strategies. This thesis has a

strong methodological component. Given the limited methodological rigor in many international segmentation studies, there is a need for developing model-based methods tailored to international segmentation problems that warrant correct and effective inference. The remaining chapters will pay particular attention to the issues of response tendencies, sampling designs, and segmentation methods as discussed in section 1.3.

The three dimensions in Figure 1.1 affect the effectiveness of international segmentation strategies in several ways. These dimensions are related to the basis of segmentation (general versus product-specific), the segmentation objects (geographic versus consumers), and segmentation methodology (heuristic versus model-based). This thesis takes a model-based approach to segmentation and considers two research directions that are related to the segmentation basis and segmentation objects in Figure 1.1.

The first direction considers the integration of product development and communication strategies in international segmentation using means-end chain theory. The consumer segmentation studies discussed in section 1.2 and depicted in Figure 1.2 either used general bases, such as values and lifestyles (Kamakura et al. 1993; Boote 1983) or product-specific bases (Moskowitz and Rabino 1994; Yavas, Verhage, and Green 1992). Means-end chain theory provides a way to link product attributes and benefits of product use (product-specific bases) to consumer values (a general basis) by postulating hierarchical links between these concepts. Using means-end chains as a basis for international segmentation has the potential to increase the actionability of targeted strategies of international product development and communication. Chapters 2, 3, and 4 develop an integrated methodology of data collection and analysis to identify such international consumer segments.

The second direction is presented in chapter 5, which focuses on the topic of improving the accessibility and cost effectiveness of international geographic segments. Along the dimension of segmentation objects, as depicted in Figure 1.1, there is a trade-off between accessibility and responsiveness of segments. Whereas consumer segments tend to be more responsive, their typical geographic

configuration is often less accessible. Especially when transportation costs strongly surpass production and marketing costs, it is important that a geographic segment defines a contiguous area as opposed to the dispersed segments that may arise from traditional cross-national segmentation approaches. The organization of this thesis is as follows.

Chapter 2 investigates the effects of international sampling designs on the estimation of mixture segmentation models. An approximate or pseudo-likelihood approach is proposed to obtain consistent estimates of segment-specific parameters when samples arise from such a complex design. The effects of ignoring the sampling design will be empirically demonstrated in the context of an international value segmentation study in which a multinomial mixture model is applied to identify segment-level value rankings.

Chapter 3 develops a measurement instrument for the identification of international segments based on means-end chains. The Association Pattern Technique (APT) is investigated as an alternative to laddering, the most popular, qualitative measurement methodology in means-end chain research. Laddering is not an appropriate measurement technique for conducting large-scale international segmentation studies. It is a quantitative technique that is time consuming, expensive, and the quality of the data may be affected by respondent fatigue and boredom. APT is a structured method that does not suffer from these limitations. It can be used in personal as well as quantitative mail interviews and is suitable for international segmentation studies. APT separately measures links between attributes and benefits and between benefits and values in two pick-any tasks. The independence of these links is crucial to the validity of APT. The assumption of independence is investigated and the convergent validity of APT and laddering is tested.

Chapter 4 proposes a methodology to identify cross-national market segments, based on means-end chain theory. The methodology offers the potential for integrating product development and communication strategies by linking product characteristics and benefits to consumer values. Building upon the estimation technique developed in chapter 2 and the APT measurement instrument validated in chapter 3, a model is developed that identifies relations between the consumer and the

product at the segment level, which in turn increases the actionability and responsiveness of the segments. Because response tendencies may hamper the identification of international segments, the model accounts for different response tendencies across and within countries. The segmentation model is applied to consumer data collected in the European Union. Additional attention is paid to model validation. A Monte Carlo study assesses the model's performance in recovering the parameters across a wide range of conditions. Relations to consumers' sociodemographics, consumption patterns, media consumption, and personality are examined. Finally, the predictive validity of the model is assessed and a comparison is made with a heuristic approach that is traditionally employed in international segmentation, i.e., K-means clustering.

Chapter 5 seeks to develop a methodology to delineate international geographic segments for companies expanding internationally. A central issue for such companies is the location of geographic segments and the design of marketing strategies that tap the needs of consumers in these segments. Whereas geography poses restrictions on the configuration of international geographic segments, its identification would ideally be based on consumer needs for formulating responsive marketing strategies. A methodology is developed that identifies alternative spatial configurations, based on the similarity of consumers' needs. The methodology integrates previous approaches by identifying geographic segments based on consumer needs, but at the same time enforces contiguity of the segments to warrant its actionability. In model specification, a hierarchical Bayes approach is taken that identifies geographic target segments obeying contiguity constraints. The methodology is applied in the international retailing domain, where international expansion is currently a significant growth strategy. A synthetic data analysis is conducted to assess the performance of the model under different experimental conditions. To assess the accessibility and responsiveness, the segments are related to relevant characteristics and are empirically compared with an unrestricted solution of the model.

Chapter 6 concludes with an integral discussion of chapters 2 to 5 and concludes with a number of limitations and issues for future research.

Chapter 2

Sampling Designs in International Segmentation¹

2.1 Introduction

In the last five to ten years, the finite mixture model approach to segmentation problems has seen an impressive upsurge in interest in the segmentation literature. In those models it is assumed that the observations come from a number of unobserved segments in the population. The normal, Poisson, Bernoulli distributions and other members of the exponential family of distributions have frequently been used to describe the observations in each segment. The mixture approach provides a statistical modeling framework for segmentation problems, in which a variety of models can be used to describe the data generation mechanism within the segments, such as regression, structural relations, hazard, and multidimensional scaling models. The mixture approach

¹ This chapter is published as Wedel, Michel, Frenkel ter Hofstede, and Jan-Benedict E.M. Steenkamp (1998). "Mixture Model Analysis of Complex Samples," *Journal of Classification*, 15 (2), 225-244.

enables the simultaneous identification and description of within-segment structure of the observations.

In finite mixture models one has typically assumed that the subjects in the sample are drawn from the population using simple random sampling. However, in practice such random samples are not necessarily desirable and seem to be the exception rather than the rule. The framework for sampling theory has been developed by Neyman (1934), who established the role of randomization as the basis for sampling strategies, and introduced the ideas of stratification and the use of unequal selection probabilities. Then a general theory of sampling was developed. In probability samples, the selection probabilities of all elements in the population are known, which allows for the projection of the estimates from the sample to the population and enables the calculation of the precision of those estimates. Next to simple random sampling, the most important probability sampling strategies are stratified sampling, cluster sampling, and two-stage sampling. We refer to such procedures as *complex* sampling procedures. Good surveys use the structure of the population and employ complex sampling designs to yield more precise estimates. The theory of probability-weighted estimation for descriptive purposes (for example estimating population totals and means) is well established (e.g., Cochran 1977, pp. 89-96). On the other hand, probability-weighted estimation for analytic purposes has received attention only fairly recently. Skinner, Holt, and Smith (1989) and Thompson (1997) provide an overview of the current status of this field of research. Survey data are frequently used for analytic purposes, such as segmentation, but in the application of statistical methods the sample design is mostly neglected.

This chapter considers the problem of how to identify unobserved segments in the population from samples arising from complex probability sampling designs. We are concerned with statistical inference about underlying segments in the population on the basis of data obtained by using a complex sample design and an international sampling design in particular. To our knowledge, this problem has not previously been dealt with in the literature. Conventional procedures for estimating segment-level parameters using mixture models are based upon the

assumption of simple random sampling. We show that if the data come from a complex probability sample, inferences on segments in the population can be made by applying approximate, or “pseudo” maximum likelihood estimators. We empirically demonstrate the effects of ignoring the sampling design in international market segmentation, using data from a stratified international value study that was recently conducted in six European countries involving nearly 4,000 consumers. We apply an ordered multinomial logit mixture model proposed by Kamakura and Mazzon (1991) to estimate cross-national value segments. This model describes the rank orders of values for segments of consumers.

2.2 Sample Design and the Mixture Approach

2.2.1 Pseudo-Maximum Likelihood Approach

The approach to deal with mixture models for complex sample designs is based on an approximate, or “pseudo” maximum likelihood (PML), estimation approach and requires the knowledge of the inclusion probabilities of each of the units selected in the sample. The development of the PML approach for mixture models is based on Skinner (1989). The methodology is originally due to Binder (1983). Assume a general complex sampling design with associated probability $p(r)$ that sample r is drawn. The design may involve combinations of sampling schemes, for example, stratified and two-stage sampling. Suppose that under the sampling design a sampling unit, denoted by n ($n = 1, \dots, N$), has an inclusion probability of $P_n = \sum_{r:n \in r} p(r)$. In general, these inclusion probabilities, $P_n(\Phi)$, will depend upon the parameters Φ of an assumed model, in our case a mixture model as detailed below. A simple random sample implies that $P_n(\Phi) = P$, for all n , since all units have the same probability of being selected into the sample, independent of the parameters of the model. Complex samples result in different inclusion probabilities for different n .

Assume that the data for subject n consist of K measurements contained in a $(K \times 1)$ vector y_n . We assume the existence of $s = 1, \dots, S$

unobserved segments, with unknown proportions π_s . Both the number of segments and their structure is unknown. Given a particular segment s , the observations are described by a probability density function $f_s(y_n|\phi_s)$, assumed to be one of the univariate exponential family: $f(y_n|\phi_s) = \exp[(y_n - b(\theta_s))/\lambda_s + c(y_n, \lambda_s)]$ (or its multivariate extension), with θ_s the natural parameter, λ_s the dispersion parameter and $b(\cdot)$ and $c(\cdot)$ functions. The exponential family includes many distributions, such as the normal, exponential, and gamma distributions for continuous variables, and the binomial, multinomial, negative binomial, and Poisson distributions for discrete variables (cf., McCullagh and Nelder 1989, p30). The common properties of these distributions enable them to be studied simultaneously, rather than as a collection of seemingly unrelated cases. This enables us to formulate the PML approach dealing with complex sampling designs for a wide class of mixture models, including those in international segmentation.

In those models, a variety of possible data generating mechanisms may be assumed, represented by a within-segment model and unknown parameter vectors ϕ_s . For example $\phi_s = (\theta_{ks})$, may consist of an intercept (or K intercepts) for segment s , leading to standard mixtures of exponential family distributions (cf., Titterton, Smith, and Makov 1985). As another example, if L independent variables, contained in a $((K+N) \times L)$ matrix X_n , are to be used to reparameterize the natural parameter as: $\theta_s = X_n \beta_s$, a mixture of generalized linear models may be assumed to underlie the data. Here, the $\Phi_s = (\beta_{ks})$ are segment-specific regression parameters (cf., Wedel and DeSarbo 1995). This class of models presents “clusterwise regression” models that enable the identification of segments of subjects and the simultaneous estimation of regression models within each segment that predict the dependent variable from a set of independent variables. Finally, in the case of mixture Multidimensional Scaling models (cf., Wedel and DeSarbo 1996) a spatial representation of a set of K stimuli in L dimensions is sought, represented by the $(K \times L)$ matrix of locations A , while simultaneously an $(L \times 1)$ preference vector γ_s (or vector of ideal points) for each segment s is identified. Here the natural parameter is reparametrized as: $\theta_s = A\gamma_s$,

and $\phi_s = (A, \gamma_s)$. (Note that in the above models any of the parameters ϕ_s may be restricted to have the same value across segments). Other models describing the within-segment structure of the observations may be employed, including hazard models (Wedel et al. 1995), or structural equation models (Jedidi, Jagpal, and DeSarbo 1997).

The unconditional distribution of the observations is:

$$(2.1) \quad f(y_n | \Phi) = \sum_{s=1}^S \pi_s f_s(y_n | \phi_s).$$

Mixture models are usually estimated by maximizing the log likelihood, providing the ML estimator of $\Phi = (\pi_s, \Phi_s)$. The ML estimator solves the equation:

$$(2.2) \quad \sum_{n=1}^N J_n(\Phi) = \sum_{n=1}^N \frac{\partial \ln f(y_n | \Phi)}{\partial \Phi} = 0.$$

The standard formulation of the log likelihood applies under simple random sampling, in which each sampling unit n receives the same weight. Ideally, a complex sample design should be incorporated into the mixture model specification. ML estimates of the model parameters can then be obtained as usual by solving the likelihood equations. Often, however, such a model-based procedure is intractable. The selection probabilities, $P_n(\Phi)$, are a function of the parameters of the model, which may render the expression for the likelihood under a particular sampling strategy highly complex. Thus, ML estimation of mixtures for complex samples is often not feasible. An alternative approach is to modify Equation (2.2) to include consistent estimates of the selection probabilities, P_n , to yield the sample estimating equations:

$$(2.3) \quad \sum_{n=1}^N \omega_n J_n(\Phi) = \sum_{n=1}^N \omega_n \frac{\partial \ln f(y_n | \Phi)}{\partial \Phi} = 0,$$

in which the weights ω_n are inversely proportional to the P_n : $\omega_n = N / (P_n \sum_{n=1}^N 1/P_n)$, so that they sum to N across the sample. Solving Equation (2.3) yields an approximate or “pseudo” maximum likelihood estimator (PML) for Φ , which is consistent. Inference proceeds

with respect to its sampling distribution over repeated samples generated from the population by the design $p(r)$ (Skinner 1989, p. 82), where Equation (2.1) is assumed to hold in the population. An advantage of the PML approach is that it provides a unifying framework for dealing with a large variety of complex sample designs, thus alleviating the problem of building mixture models that accommodate the specific sample design and the structure of the observations on a case by case basis. We note that Equation (2.3) can be derived from the formulation of the exponential family provided by Fahrmeier and Tutz (1991, p19). Complex sampling designs for which all selection probabilities are equal are called “self weighting.” For such sampling designs the ML and PML estimators coincide. Neglecting the sampling design for samples that are not self-weighting leads to inconsistent estimators. In the next section we provide the weights for a few common complex sampling designs.

2.2.2 Stratified Sampling

In stratified sampling it is assumed that the population is grouped into $g = 1, \dots, G$ strata. Samples collected for international market research and segmentation are stratified by country with strata represented by countries. Whereas population sizes vary highly across countries, national sample sizes are usually taken to be of similar magnitude, at least not being proportional to population sizes. Within stratum g , N_g subjects are sampled. Let $N_g^{(p)}$ indicate the corresponding number of units in stratum g in the population. The sample size is $N = \sum_{g=1}^G N_g$ and the population size is $N^{(p)} = \sum_{g=1}^G N_g^{(p)}$. A mixture model is to be applied to the N ($K \times 1$) observation vectors y_n . The appropriate PML estimates of the parameters are weighted estimates obtained from Equation (2.3), in which case the selection probabilities equal: $P_n = N_{g(n)} / N_{g(n)}^{(p)}$, where $g(n)$ is the stratum from which subject n comes. If the sampled fraction is constant across strata $P_n = P$ the sample is self-weighting and the ML and PML estimators coincide.

2.2.3 Cluster Sampling

In cluster sampling a sample of primary units m (often referred to as clusters) is drawn from a population of $M^{(p)}$ units. Within each primary unit observations on all secondary units $n = 1, \dots, N_m$ are obtained, where the sample size is $N = \sum_{m=1}^M N_m$ and the size of the population is

$N^{(p)} = \sum_{m=1}^{M^{(p)}} N_m^{(p)}$. Since cluster samples contain all secondary units in the selected primary units, the inclusion probabilities of the secondary units equal those of the primary units ($P_{nm} = P_m$). For example, when a simple random sample is drawn from primary units all having the same size, then $P_{nm} = 1/N$. The PML and ML estimators coincide in this case. If the selection probabilities of the primary units differ, the sampling design is not self-weighting and should be accommodated in estimation. See Cochran (1977, p. 233-270) for further results.

2.2.4 Two-Stage Sampling

Like in cluster sampling, in two-stage sampling a sample of size M is drawn from the primary units in the population. However, in two-stage sampling, a *sample* of secondary units of size N_m is drawn from each primary unit that has been drawn in the first stage. If, for example, (a) the primary units in the population have the same size, and (b) the primary and the secondary units are selected by simple random sampling in a constant fraction, the sample is self-weighting: $P_{nm} = 1/N$. If the primary and secondary units are selected by simple random sampling and the sizes of the primary units differ, the selection probabilities are: $P_{nm} = (M \cdot N_m) / (M^{(p)} \cdot N^{(p)})$. If the sampling fraction within each primary unit is constant, the sampling strategy is again self-weighting. Further results can be derived from Cochran (1977, p. 274-324).

2.3 Statistical Inference

2.3.1 Estimation

The standard method for estimating mixture models is to maximize the likelihood function (e.g., McLachlan and Basford 1988, Titterington, Smith, and Makov 1985). The likelihood of finite mixtures can be maximized by solving Equation (2.2) using either standard numerical optimization routines such as Newton-Raphson or Conjugate Gradients (cf., Scales 1985, p. 61-83), or by using the Expectation-Maximization (EM) algorithm (Dempster, Laird, and Rubin 1977). PML estimates can similarly be obtained by solving Equation (2.3) using either numerical optimization or the EM algorithm. Numerical methods require relatively few iterations to converge, but convergence is not ensured. In the EM algorithm convergence is ensured, but the algorithm may require many iterations. Both EM and numerical estimation methods may converge to local optima. A solution to this problem is to have the algorithms started from several sets of starting values for the parameters and to inspect the estimates for local optima. It is at present not clear which of the two methods is to be preferred in general.

2.3.2 Asymptotic Standard Errors of the Estimates

Under typical regularity conditions, the ML estimators are efficient and asymptotically normal. A consistent estimator of the asymptotic covariance matrix of the estimates is provided by the inverse of the observed Fisher information matrix (e.g., Titterington, Smith, and Makov 1985). The PML estimator obtained from Equation (2.3), however, is not efficient. The reason is that introducing as weights the selection probabilities decreases the efficiency of the estimator. (This points to the advantages of using self-weighting samples.) A consistent and robust estimator of the asymptotic variance of $\hat{\Phi}_{PML}$ is provided by White (1982), and Royall (1986):

$$(2.4) \quad \Sigma(\hat{\Phi}_{PML}) = I(\hat{\Phi}_{PML})^{-1} V(\hat{\Phi}_{PML}) I(\hat{\Phi}_{PML})^{-1},$$

where $I(\hat{\Phi}_{PML})$ is the observed information matrix, and $V(\hat{\Phi}_{PML})$ is based on the cross product of the vector of first derivatives, $J_n(\hat{\Phi}_{PML})$, summed across strata and/or primary sampling units (see White 1982, and Royall 1986 for details). In general, the sample size needs to be fairly large in mixture model analysis for the asymptotic approximations to the standard errors in ML and PML estimation to be adequate.

2.3.3 Criteria for Selecting the Number of Segments

Information criteria are commonly used to identify the most appropriate mixture model. They impose a penalty, involving the number of parameters estimated (Q), upon minus two times the log likelihood:

$$(2.5) \quad C_{ML}(S) = -2 \ln L(\hat{\Phi}_{ML} | S) + Qd.$$

Here, d is some constant. Akaike's Information Criterion, AIC (Akaike 1974), arises when $d = 2$. For the Bayesian Information Criterion, BIC (Schwartz 1978), $d = \ln(N)$ and for the Consistent Akaike's Information Criterion, CAIC (Bozdogan 1987), $d = \ln(N+1)$. For ICOMP (Bozdogan 1990): $d = -\ln \text{tr}[I(\hat{\Phi}_{PML})^{-1}] - \frac{1}{Q} \ln \det[I(\hat{\Phi}_{PML})^{-1}]$. ICOMP penalizes the likelihood more when the model becomes less well identified due to an increasing number of parameters; its value depends on the scaling of the variable in question.

Assume random sampling, the standard regularity conditions, a true model specified by H_0 and an alternative specified by H_1 . Then, for $N \rightarrow \infty$ the probability that the $C_{ML}(S)$ statistics select H_1 is: $P_{ML}(\chi^2_{Q_0-Q_1} > Q_0 d_0 - Q_1 d_1)$, which is larger than zero for AIC, but equals zero for the other statistics (Bozdogan 1987). Thus, the AIC statistic tends to overstate the number of segments asymptotically, but the other statistics are consistent in indicating the dimensionality of the model.

However, problems arise when using information statistics to select the number of segments in mixture model analysis of complex samples. First, if ML is applied ignoring a complex sampling design, the

ML estimator, $\hat{\Phi}_{ML}$, is inconsistent. Therefore the $C_{ML}(S)$ statistics are not dimension consistent and will overstate the number of segments. Second, the difference in the log likelihoods for testing S against $S+1$ segments may not be asymptotically distributed as Chi-square, because of unidentifiability of some mixtures when one of the π_s equals zero (Titterington, Smith, and Makov 1985, p4) and because the S -segment solution corresponds to a boundary of the parameter space for the $(S+1)$ -segment solution, which may violate the required regularity conditions (cf., Aitkin and Rubin 1985).

We use the information statistics based on the log pseudo likelihood:

$$(2.6) \quad C_{PML}(S) = -2 \ln PL(\hat{\Phi}_{PML} | S) + Qd,$$

with:

$$(2.7) \quad \ln PL(\Phi | S) = \sum_{n=1}^N \omega_n \ln \sum_{s=1}^S \pi_s f_s(y_n | \phi_s).$$

Since the PML estimator $\hat{\Phi}_{PML}$ is consistent (Skinner 1989, p.82), the type I error for the $C_{PML}(S)$ statistics is zero asymptotically (except for AIC), so that they are dimension consistent. This fact motivates the use of the $C_{PML}(S)$ over the $C_{ML}(S)$ statistics. However, the problems of the statistics in testing S versus $S+1$ segments carry over from the ML situation, so that the asymptotic approximations to the difference in the pseudo log likelihoods are not valid. Therefore, our approach to determine S involves using the minimum $C_{PML}(S)$ rule as a heuristic.

2.4 Illustration

2.4.1 Background

We illustrate the PML procedure in the context of cross-national value segmentation. In the social sciences, there is a growing interest in the concept of human values (Schwartz 1992). Values underlie a large and important part of human cognition and behavior. An examination of

values provides both an overall picture of a central cognitive structure of the individual, as well as a means of linking central beliefs to attitudes.

A value is defined as an enduring belief that a specific state of existence or mode of conduct is preferred for living one's life (Rokeach 1973). Rokeach (1973) proposed the concept of the value system in which each value is ordered in priority relative to other values, to account for the fact that subjects can have several, sometimes conflicting values. The value system provides a comprehensive understanding of the motivational forces driving an individual's beliefs, attitudes, and behavior. The advantages of analyzing individuals at the level of the complete value system level, rather than at the level of individual values, has been recently demonstrated by Kamakura and Mazzon (1991) and Kamakura and Novak (1992).

Values provide potentially powerful explanations of human behavior because they serve as standards of conduct, tend to be limited in number, universal across cultures, and temporally stable (Kamakura and Mazzon 1991, Rokeach 1973, Schwartz 1992). Consequently, it is not surprising that the concepts of human values and value systems have been widely used by social scientists to explain a variety of attitudinal and behavioral outcomes including charity contributions (Manzer and Miller 1978), mass media usage (Rokeach and Ball-Rokeach 1989), religious behavior (Feather 1984), societal cultural orientation (Schwartz 1992), participation in civil rights activities (Rokeach 1973), cheating on examinations (Henshel 1971), job performance appraisals (Jolly, Reynolds, and Slocum 1988), drug addiction (Toler 1975), intergroup relations (Schwartz 1994), political orientation (Schwartz 1994), organization behavior (Clare and Sandford 1979), innovativeness (Steenkamp, ter Hofstede, and Wedel 1999), and purchase of organic foods (Grunert and Juhl 1995).

One area in which values and value systems have found broad application is market segmentation (e.g., Kamakura and Mazzon 1991; Kahle, Beatty, and Homer 1986; Novak and MacEvoy 1990). Values are a useful basis for segmenting consumers because values are more closely related to behavior than are personality traits, and they are less numerous, more central, and more immediately related to motivations than attitudes.

The universality, centrality, and stability of values render values not only a particularly suitable basis for segmentation *within* a country but also for *cross-national* market segmentation. From a substantive point of view, the issue of cross-national segmentation relates to one of the key international marketing issues of the 1990's, viz. the extent to which firms can design pan-regional (or global) marketing strategies. If no cross-national segments exist, a pan-regional strategy is unlikely to be effective because it may fail to appeal to prospective buyers in different countries (Douglas and Craig 1992; Jain 1989).

2.4.2 Data

The PML procedure was used in a stratified European value segmentation study. The effects of taking the sampling design into account are demonstrated by comparing the PML estimates to the ML estimates. The purpose of the study is to investigate the existence of pan-European value segments, i.e., segments that transcend national borders. The data were collected in 1996 in six West European countries by an international marketing research agency. The sample was stratified by country. From each country a sample of approximately the same size was drawn, although the population sizes in the countries differ substantially. This procedure is standard in international marketing research. The countries were (sample sizes/ weights in parentheses): Belgium (648/ .231), France (694/ 1.259), Germany (673/ 2.020), Great Britain (623/ 1.356), the Netherlands (646/ .380), and Spain (616/ .690).

Several instruments have been developed to measure a person's value system. One of the most popular instruments is Kahle's (1986) "List of Values" (LOV). The LOV method was also used in the present study to assess value systems of respondents. The LOV typology is related to social distinction theory. It distinguishes between external and internal values and deals with the importance of interpersonal relations, and personal and a-personal factors in value fulfillment. The LOV instrument is composed of nine values. These were ranked in a paper and pencil task by the respondents. Back-translation methods were used to ensure a similar content of the statements in the languages involved (Brislin 1970). Thus, the data consist of the order of the nine values, for

each of 3900 respondents, where the lower the rank number the more important a value is for a person.

2.4.3 The International Value Segmentation Model

To identify latent value segments, a mixture of rank order multinomial logit models is used, proposed by Kamakura and Mazzon (1991). The model (a) is a mixture model extension of Thurstone's (1927) model based on his law of comparative judgment, in which it is assumed that the observed value rankings are error-perturbed observations of the unobservable value utilities of each individual, (b) is based on utility maximization theory, and (c) identifies segments on the basis of the entire value ranking provided by the subjects. The model is developed as follows. Let y_{nt} denote the observed value label on rank order position t , $t = 1, \dots, T$ ($T=9$ for LOV) for subject n . For example, $y_{n1} = 6$ means that subject n orders value number 6 (i.e., excitement) on position 1. Assume the existence of S unobserved value segments. An individual n , belonging to segment s has unobserved utility, $u_{nt|s}$ for the value at position t ($t = 1, \dots, T$), given segments. According to Thurstone's (1927) theory, these utilities are assumed to be random, i.e., $u_{nt|s} = v_{t|s} + \varepsilon_{nts}$. Thus, individuals in segment s share a single underlying value system, represented by the (deterministic) set of utilities $v_{t|s}$. The error terms, ε_{nts} , in the utility are assumed to be i.i.d. type I extreme value distributed with unit variance, i.e., their cumulative distribution function is $F(\varepsilon_{nts}) = \exp(-\exp(-\varepsilon_{nts}))$ (Amemiya 1985, pp. 296-299). Given segment s , the probability that the utility of this value to subject n is larger than the utilities of values ranked higher is given by:

$$\begin{aligned}
 (2.8) \quad P(y_{nt}) &= P(u_{nt|s} > u_{nq|s}, \forall q, q > t) \\
 &= P(\varepsilon_{nqs} \leq v_{t|s} - v_{q|s} + \varepsilon_{nts}, \forall q, q > t)
 \end{aligned}$$

Thus, Equation (2.8) defines a joint cumulative distribution function on the ε_{nqs} , which yields (McFadden 1974, p. 111):

$$\begin{aligned}
 P(y_{nt} | v_{1|s}, \dots, v_{T|s}) &= \int_{e_t=-\infty}^{\infty} \exp(-e_t) \prod_{q:q \geq t} \exp(-\exp(v_{q|s} - v_{t|s} - e_t)) de_t. \\
 (2.9) \quad &= \frac{\exp(v_{t|s})}{\sum_{q:q \geq t} \exp(v_{q|s})}.
 \end{aligned}$$

Equation (2.9) is not identified, however, because a constant can be added to the $v_{t|s}$ without affecting the result. Therefore, one of the $v_{t|s}$ is set to zero. We choose to set the utility of LOV-9 to zero for all segments in order to enable comparison of utility estimates across segments, but other identifying constraints (e.g., the sum or mean of the utilities can be constrained to zero) may be imposed without changing the results. The log-pseudo likelihood is:

$$(2.10) \quad \ln PL(\Phi | y) = \sum_{n=1}^N \omega_n \ln \sum_{s=1}^S \pi_s \prod_{t=1}^T P(y_{nt} | v_{1|s}, \dots, v_{T|s}),$$

with $\omega_n = (N_{g(n)}^{(p)} \cdot N) / (N_{g(n)} \cdot N^{(p)})$. The log likelihood is obtained from Equation (2.10) by setting $\omega_n = 1$ for all n . Note that the sample size is large ($N=3900$) so that the asymptotic approximations are likely to be adequate. The models are estimated using a conjugate gradient search to solve for Φ , using the numerical optimization procedures implemented in the GAUSS package (Aptech 1995).

2.4.4 Comparison of ML and PML Results

We estimate the model on the LOV data, using both the ML and the PML approach. Starting with $S = 1$, S is increased until the AIC, CAIC, BIC and/or ICOMP show a minimum. To overcome problems of local optima, each model is estimated from 10 sets of random starting values. Table 2.1 shows the values of the log likelihood and the model selection criteria for ML estimation. Table 2.2 shows the statistics for PML estimation. The solutions with the highest likelihood values out of the 10 repeated analyses are reported. The likelihoods of those 10 solutions were always within 1.5% of the maximal value, so that the problems of local optima seem to be minor. For ML estimation, BIC, CAIC, and ICOMP, all reach

a minimum at $S=7$ (see Table 2.1). Since AIC tends to overstate the number of segments, the seven-segment ML solution appears to be a reasonable representation of the data. For PML estimation, on the other hand, BIC, CAIC, and ICOMP all indicate that the $S=5$ solution is most appropriate. Thus, the first important finding is that not accounting for stratification in ML estimation may yield a too large number of segments. Accounting for the stratified sampling procedure through PML estimation results in more parsimonious models (the $S=5$ solution has 18 parameters less than the $S=7$ solution). BIC, CAIC and ICOMP identify the same number of segments for ML and PML estimation, while AIC tends to indicate an excessive number of segments.

Table 2.1

Model selection criteria for ML estimation[†]

<i>S</i>	<i>-lnL</i>	<i>AIC</i>	<i>BIC</i>	<i>CAIC</i>	<i>ICOMP</i>
1	45458.56	90933.13	90983.28	90991.54	90937.34
2	44892.24	89818.50	89925.06	89933.33	89842.63
3	44734.87	89521.75	89684.74	89693.01	89574.28
4	44530.97	89131.94	89351.35	89359.62	89213.08
5	44454.78	88997.56	89273.39	89281.66	89109.48
6	44397.70	88901.41	89233.65	89241.92	89056.67
7	44349.42	88822.95	89211.51	89219.78	89034.50
8	44336.46	88814.92	89260.00	89268.27	89322.11

[†] Boldface type indicates the minimum value across $S=1$ to $S=8$.

Table 2.2

Model selection criteria for PML estimation[†]

<i>S</i>	<i>-lnPL</i>	<i>AIC</i>	<i>BIC</i>	<i>CAIC</i>	<i>ICOMP</i>
1	45458.56	90933.13	90983.28	90991.54	90937.34
2	44892.25	89818.49	89925.06	89933.33	89842.63
3	44734.88	89521.75	89684.74	89693.01	89576.66
4	44530.97	89131.94	89351.34	89359.61	89211.57
5	44471.89	89031.78	89307.61	89315.88	89163.49
6	44452.25	89010.50	89342.74	89351.01	89601.36

[†] Boldface type indicates the minimum value across $S = 1$ to $S = 6$.

Table 2.3

Prior probabilities for S=5 PML and ML solutions, aggregate and per country

<i>Segment</i>	<i>BE</i>	<i>GE</i>	<i>GB</i>	<i>FR</i>	<i>NL</i>	<i>SP</i>	<i>Aggregate</i>
PML 1	.387	.246	.278	.542	.385	.167	.338
PML 2	.038	.325	.115	.020	.047	.022	.095
PML 3	.202	.068	.192	.074	.213	.212	.158
PML 4	.004	.138	.015	.002	.004	.009	.029
PML 5	.370	.223	.401	.363	.351	.591	.380
ML 1	.326	.257	.231	.435	.350	.123	.290
ML 2	.034	.431	.122	.020	.044	.035	.116
ML 3	.236	.102	.240	.090	.246	.251	.191
ML 4	.139	.053	.116	.198	.091	.125	.121
ML 5	.266	.157	.291	.258	.269	.467	.282

Next, we inspect the ML and PML estimates. We focus on the $S=5$ solution, since that is the best approximation to the data, when the appropriate weighting is applied. The solutions are optimally matched according to the correlations between the utility estimates. The prior probabilities (or segment proportions) are presented in Table 2.3, which shows a number of differences between the $S=5$ ML and PML estimates. First, the aggregate segment proportions (π_s) are quite different between ML and PML. For PML, Segment 4 is much smaller (3%) and Segment 5 much larger (38%) than the corresponding segments in the ML solution (12% and 28%). It appears that ML has a tendency to identify segments with more equal sizes. The comparison shows that the ML estimates are severely biased because of their indifference to the sampling design. Furthermore, the PML and ML solutions show substantial differences in the estimated segment proportions per country. In particular, PML Segment 4 has particularly low proportions of Belgium, Great Britain, France, the Netherlands and Spain (.4%, 2%, .2%, .4% and .9%, respectively), while for ML those proportions are much higher (14%, 12%, 20%, 9% and 12%). In Germany, however, the proportion in Segment 4 is much lower for ML than for PML (5% versus 14%). The

proportions of Great Britain and Spain in PML Segment 5 (40% and 59%) are substantially larger than those in ML Segment 5 (29% and 47%).

Table 2.4

Value utilities in each of five segments for S = 5 PML and ML solutions

<i>Segment</i>	<i>LOV-1</i>	<i>LOV-2</i>	<i>LOV-3</i>	<i>LOV-4</i>	<i>LOV-5</i>	<i>LOV-6</i>	<i>LOV-7</i>	<i>LOV-8</i>
PML 1	2.324	1.140	-.350	-.119	-.980	-1.228	-.301	-.163
PML 2	1.140	1.016	-1.569	-1.193	1.249	-2.276	-1.158	1.093
PML 3	-.621	1.903	-.849	.464	.071	-2.35	.114	.133
PML 4	.874	-2.597	-2.158	-2.425	.492	-3.27	-1.525	.912
PML 5	-1.197	-.664	-1.077	-.769	-1.988	-2.928	-.808	-.356
ML 1	2.746	1.365	-.532	-.098	-.868	-1.264	-.377	.009
ML 2	.786	-.354	-1.394	-1.301	.895	-2.231	-1.080	.994
ML 3	-.579	1.808	-.851	.292	.009	-2.35	-.005	.137
ML 4	.388	.093	.383	-.567	-1.690	-1.723	-.084	-.876
ML 5	-1.573	-.929	-1.432	-.866	-2.310	-3.345	-1.051	-.265

LOV-1 = "Fun and enjoyment in life," LOV-2 = "Warm relationships with others," LOV-3 = "Self-fulfillment," LOV-4 = "Being well respected," LOV-5 = "Sense of belonging," LOV-6 = "Excitement," LOV-7 = "A sense of accomplishment," LOV-8 = "Security," LOV-9 = "Self respect."

Table 2.4 depicts the value-utility estimates obtained with ML and PML for the $S = 5$ solution. The solutions are optimally matched based on the correlations between the utilities. There is a high correlation of the ML and PML estimates for Segment 1 (1.00), Segment 2 (.95), Segment 3 (1.00), and Segment 5 (.99). ML Segment 4 correlates with PML Segments 1 (.72) and 5 (.62), but the correlation with PML Segment 4 is low (-.02). The value rankings predicted for Segment 4 are very different between the ML and PML procedures, with ML predicting high ranks for warm relationships with others (LOV-2) and self-fulfillment (LOV-3), while PML predicts low rankings for those values. ML on the other hand predicts low rankings for sense of belonging (LOV-5), and security (LOV-8) while PML predicts high rankings for those values. The estimates for the values in the other segments also show some

differences, however. For example, in Segment 2, the PML procedure predicts a higher value ranking for warm relationships with others (LOV-2) as compared to ML. In general, ML and PML yield different results and the ML procedure seems not to recover Segment 4 in particular.

Finally, we show the extent to which the actual classification of the subjects in the sample corresponds between ML and PML. For that purpose, we assigned all subjects to that segment for which the posterior probability of classification was largest (if Φ were known this would be the optimal Bayes rule of classification, McLachlan and Basford 1988, p. 11). Then, the memberships in the ML and PML segments were cross-tabulated as shown in Table 2.5 for the $S=5$ solution.

Table 2.5

Cross-classification of membership between $S=5$ PML and ML solutions

	<i>ML 1</i>	<i>ML 2</i>	<i>ML 3</i>	<i>ML 4</i>	<i>ML 5</i>	<i>Total</i>
PML 1	1197	34	28	157	4	1420
PML 2	77	276	33	0	0	386
PML 3	0	0	529	9	1	539
PML 4	0	109	0	0	0	109
PML 5	11	48	122	188	1077	1446
Total	1285	467	712	354	1082	3900
%	93	59	74	0	99	75

Table 2.5 shows that the classification of subjects into the five value-segments on the basis of the posteriors is quite different for ML and PML. The assignment of subjects to Segments 1, 3, and 5 is quite similar, ML predicting 93%, 74%, and 100% of the memberships correctly relative to PML. However, Segments 2 and 4 are not recovered very well by ML. None of the subjects in ML Segment 4 is assigned to the same segment with PML. This explains the very large differences in value utility estimates for Segment 4 reported above. This result explains the very large differences in value utility estimates for Segment 4 reported above. The results of the matching of the segments on the basis of the

cross-classification in Table 2.5 corresponds to the matching on the basis of the correlations of the utility estimates in Table 2.4.

2.4.5 Substantive Interpretation of the Five-segment PML Solution

Given the fact that the $S=5$ PML model presents the proper description of the structure of international values segments, we provide a more detailed interpretation of the results. Inspecting the ranks of the estimated value utilities among segments yields valuable insights. The orders of the values in the five segments are: *Segment 1*: 1, 2, 9, 4, 8, 7, 3, 5, 6; *Segment 2*: 5, 1, 8, 2, 9, 7, 4, 3, 6; *Segment 3*: 2, 4, 8, 7, 5, 9, 1, 3, 6; *Segment 4*: 8, 1, 5, 9, 7, 3, 4, 2, 6; *Segment 5*: 9, 8, 2, 4, 7, 3, 1, 5, 6. Thus, LOV-6 "Excitement" always has lowest utility. Segments are clearly distinguished by the highest ranked values LOV-1 "Fun and enjoyment in life," LOV-2 "Warm relation with others," LOV-5 "Sense of belonging," LOV-8 "Security," and LOV-9 "Self respect." Note though that the entire value system and especially the most important values, rather than only the top value should be taken into account to describe segments fully (Kamakura and Novak 1992).

Segment 1 assigns most priority to fun and enjoyment in life, with the second highest ranking for warm relationships with others. Segment 1 consists of hedonistic individuals and is a truly cross-national segment, since each country has a substantial number of consumers belonging to this segment.

Segment 2 is characterized by four values of nearly equal importance, viz., sense of belonging, warm relationships with others, security, and fun and enjoyment in life. This segment is more socially oriented and seems to derive its fun and enjoyment through social interaction. Segment 2 has substantial memberships in Germany (33%) and, to a lesser extent, Great Britain (12%). The size of this segment in the other countries is always below 5%. Segment 3 assigns top priority to warm relationships with others while assigning some of the lowest importances to fun and enjoyment, self-fulfillment, and excitement. Segment 3 is reasonably cross-national in that it has fairly high memberships in each country, with the possible exception of Germany and France where memberships are somewhat lower.

Segment 4 exhibits partially the same profile as Segment 2, in that sense of belonging, security, and fun and enjoyment are the three most important values. However, the two segments differ especially in the importance attached to warm relationships with others. While Segment 2 assigns a high priority to this value, it is the second least important value in Segment 4. Nearly 14% of the Germans belong to Segment 4 and its size in the other countries does not exceed 2% of the respective populations. Thus, this segment is essentially a German segment. Finally, Segment 5 is the only segment that gives top priority to self-respect. Just as with Segment 1, this segment is cross-national; each country has a substantial membership in it.

Thus, two segments were identified that are truly cross-national: Segment 1 and Segment 5. These two segments, comprising 72% of the population in the six countries involved offer the best potential for pan-European marketing strategies (cf., Guido 1991, Jain 1989). On the other hand, a country-specific approach is likely to be more effective for Segments 2 and 4. Once the value motivations that drive each segment are delineated, it becomes easier to make predictions regarding the pattern of beliefs, attitudes, and behavior expected from each segment (Kamakura and Novak 1992). Values are especially suitable for segmenting the market when behavioral patterns across one or more product categories are the focal point of interest (Kamakura et al. 1993). From a managerial perspective, patterns of behavior are of special interest when the firm offers multiple product categories (e.g., soaps, cosmetics, etc.), which relate to broader behavioral domains (e.g., personal care). The value segmentation approach takes communalities in behavioral patterns into account, and thus leads to efficiencies in the design of marketing research and marketing strategy.

2.5 Conclusions

The purpose of this chapter is to investigate the effects of sample design on maximum likelihood estimation of mixture segmentation models and to propose a framework for accommodating those effects. To our

knowledge, problems resulting from complex sample designs have not previously been raised in the segmentation literature. The contribution of this study is to show how a relatively simple pseudo maximum likelihood procedure can be used to accommodate complex sample design in the estimation and evaluation of finite mixtures of distributions in the exponential family. In the empirical application, we demonstrated the effects of the PML approach for international segmentation. The results suggest that the estimates of the number of segments, segment proportions and segment-level parameters may be severely biased when using ML instead of PML. In Chapter 4 the PML approach is applied in an international segmentation study based on relations between consumer values and product attributes. First, in the next chapter, we will propose and validate a measurement technique that measures such relations.

Chapter 3

An Investigation into the Association Pattern Technique¹

3.1 Introduction

Means-end chain theory has proven to be useful in understanding consumer behavior (Peter and Olson 1987; Valette-Florence and Rapacchi 1991; Aurifeille and Valette-Florence 1995; Pieters, Baumgartner, and Allen 1995). Means-end chain theory builds upon work of psychologists (Tolman 1932) and economists (Abbott 1955) who have long recognized that consumers do not buy products for the product's sake, but for what the product can do for them. The theory relates the product to the consumer by positing a hierarchical cognitive

¹ This chapter is published as ter Hofstede, Frenkel, Anke Audenaert, Jan-Benedict E.M. Steenkamp, and Michel Wedel (1998), "An Investigation into the Association Pattern Technique as a Quantitative Approach to Measuring Means-End Chains," *International Journal of Research in Marketing*, 15 (January), 37-50.

structure involving links between attributes of the product, consequences or benefits of product use, and values of consumers. These concepts constitute the *content* of consumer knowledge, whereas the links between the concepts constitute the *structure*. Means-end chain theory, thus, is concerned with the content and structure of consumer knowledge. It extends previous work in psychology and economics in that it is more specific in what the content and structure of consumer knowledge is, and how it can be assessed and applied to practical marketing problems, such as advertising (Reynolds and Craddock 1988) and product development (Reynolds and Gutman 1988).

In means-end chain theory, products are seen as means through which consumers obtain valued ends. According to this theory, consumers choose products because they believe that the specific attributes of the product can help them to achieve desired values through the consequences or benefits of product-use (Reynolds and Gutman 1984). Attributes are the concrete, tangible characteristics of the product. Benefits refer to what the product does or provides to the consumer at the functional or psychosocial level. Values are intangible, higher-order outcomes or ends, being cognitive representations of consumers' most basic and fundamental needs and goals. The three levels: attributes, benefits and values are postulated to be hierarchically structured in that attributes lead to benefits, which produce value satisfaction. Although most models of cognitive structure specify some kind of hierarchical structure (e.g., Bettman 1974; Schank 1982), the defining feature of means-end chain theory is the specification of the exact links between the levels in the hierarchy, constituting the structure of consumer knowledge.

Despite the attention means-end chain theory has generated, there are still a number of issues that need to be addressed (cf., Grunert and Grunert 1995). One particularly important issue concerns the methodology to measure consumers' means-end chains. Laddering is by far the most popular methodology (Reynolds and Gutman 1988; Claeys, Swinnen, and Vanden Abeele 1995). It is a qualitative interviewing technique and has been applied successfully in academic and applied research, but it is not without its limitations. Laddering interviews are time-consuming, costly, and require highly trained interviewers.

Consequently, it is difficult, if not impossible, to use them to obtain the large-scale representative samples that are typically required for international segmentation studies.

In this chapter we describe and investigate a quantitative methodology to measure means-end chains, which we will denote as the Association Pattern Technique (APT). APT uses a fixed-format, and measures the links between attributes and benefits, and the links between benefits and values separately. The technique is much cheaper and faster than laddering, can be used in mail questionnaires, and allows the researcher to collect data among a representative sample of consumers. Thus, APT may be seen as an attractive option to measure means-end chains when large samples are needed, such as in comprehensive international surveys. The purpose of this chapter is to empirically explore the validity of APT. More specifically, we investigate the crucial assumption made in APT, viz. the independence of the attribute-benefit (AB) links and benefit-value (BV) links. In addition, the convergent validity between APT and laddering in terms of content and structure of the means-end chains is investigated.

3.2 Techniques for Measuring Means-end Chains

3.2.1 Laddering

Several approaches have been suggested in the literature to reveal means-end structures prevalent among consumers. Among these approaches is 'laddering' the most widely applied technique (Reynolds and Gutman 1988). In fact, sometimes laddering is equated with means-end chain theory, although the theory should be considered separate from the methodology. In the laddering procedure, three steps are distinguished: (1) elicitation of salient attributes, (2) the depth-interview, and (3) analysis of the results. In the first phase, the attribute-elicitation phase, consumers are questioned about the attributes used to compare and evaluate products. Several techniques may be used to elicit the important

attributes, including the repertory grid, stimulus grouping, and direct elicitation.

In the second step, the most important attributes identified in the first phase are used as a starting point for the depth interview. The consumer is continuously probed with some form of the question “why is that important to you?” This way of questioning forces the subject up the ‘ladder’ of abstraction, until s/he cannot go further. The end need not always be at the value level. The result is a sequence of concepts, which are called ladders.

In the third step, the idiosyncratic concepts resulting from the laddering interviews are categorized into a smaller number of categories. The links between the (categorized) concepts are represented in the so-called implication matrix. This is an asymmetrical dominance matrix in which the concepts constitute both the rows and the columns. The cell entries indicate the frequencies, across all subjects in the sample, with which an attribute, benefit or value (the row element) leads directly or indirectly (through one or more other concepts) to another attribute, benefit or value (the column element). The implication matrix preserves information about the sequence of concepts in the means-end structure, but discards differences between ladders of the same or different individuals. From the implication matrix, the so-called hierarchical value map is constructed, depicting the content and structure of consumer knowledge in a graphical way. The map gives an aggregate network representation of the means-end chains for the product in question (Reynolds and Gutman 1988).

In the classical laddering procedure described above, the natural flow of speech of the consumer is restricted as little as possible. This kind of laddering is referred to as ‘soft laddering’. On the other hand, some researchers used the ‘hard laddering’ approach, allowing less freedom in the answers of consumers and forcing consumers to follow one ladder at a time, in which each subsequent answer is on a higher level of abstraction (Grunert and Grunert 1995). Examples of this ‘hard laddering’ approach are the self-administered paper and pencil method (Walker and Olson 1991; Young and Feigin 1975) and computerized data collection methods.

Laddering has served as a very useful qualitative technique to reveal means-end chains. However, it also has its limitations. Because the laddering interview is time consuming and must be carried out by trained interviewers, it is an expensive data collection technique. Moreover, it places a serious burden on respondents, and the quality of the data may be affected by respondent fatigue and boredom (cf., Steenkamp and Van Trijp 1996). In sum, laddering is not suitable as an instrument to be used in large representative international samples, nor was it intended to be used in this context.

In reaction to these limitations there have been several attempts to quantify means-end chains in large-scale studies using other methods. Vanden Abeele (1992) used a technique where respondents evaluate complete ladders. These ladders are explicitly described in laddering statements and rated on their credibility. The method may be used in questionnaires for large studies and provides quantitative information on ladders. However, the ladders are defined a priori and the number of ladders that may be processed is limited. Grunert (1997) used an extended form of conjoint analysis to quantify means-end chains of quality perceptions. Respondents rate a few product profiles on several characteristics. Structural equations models were used to estimate relations between characteristics at the aggregate level. Still the question remains as to how the estimates relate to the actual strengths of links. In this chapter we propose the Association Pattern Technique (APT) as a quantitatively oriented technique to assess means-end chains.

3.2.2 The Association Pattern Technique

The Association Pattern Technique is inspired by Gutman (1982). Gutman proposed that, for measurement purposes, the means-end chain can be conceived as a series of connected matrices. In APT an AB-matrix (attribute-benefit matrix) and a BV-matrix (benefit-value matrix) are distinguished. Figure 3.1 gives an example of AB- and BV-matrices that were used in our empirical study (see below). In the AB-matrix, the a priori defined attributes and benefits are listed in the columns and rows, respectively, resulting in a table of all combinations of attributes and benefits. Similarly, the BV-matrix includes all possible combinations of

the benefits and values. For each column in the AB-matrix (BV-matrix), respondents indicate to which benefits (values) that attribute (benefit) is perceived to lead. This results in a data set of binary observations. Note that, in contrast to laddering, the attributes, benefits, and values are to be provided by the researcher. Since these concepts need to be relevant and need to cover the range of concepts that constitute the content of means-end chains, pretesting is inevitable when secondary sources are lacking.

Given the limitations of laddering, the Association Pattern Technique may serve as a useful supplement to laddering for measuring means-end chains. APT addresses the above mentioned limitations of laddering. The data-collection process is structured and it can be used in large-scale studies with appropriate control over representativeness of the sample on important characteristics. Obvious advantages accrue, for example in international marketing research and segmentation, where sample sizes are large and representativeness is important to obtain accurate market forecasts. Furthermore, the analysis of the data is simple, especially because of the standardization of the concepts used in the matrices. This renders content analysis unnecessary.

3.3 Key Research Issues

Given these advantages of APT, the legitimacy of APT as a method to uncover means-end chains becomes an important issue. We make a distinction between two key research issues: the assumption of independence of the AB- and BV-links, and the convergent validity between APT and laddering. Below, we will elaborate further on these two issues.

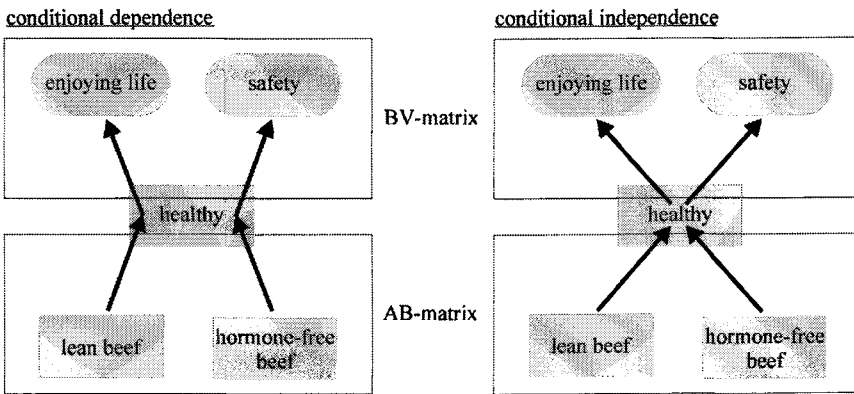
3.3.1 Assumption of Conditional Independence.

The basic assumption underlying APT is that AB- and BV-links are independent. This is a consequence of the separation of means-end chains in AB- and BV-matrices. The APT procedure implicitly assumes that the link of a benefit to a certain value in the BV-matrix is independent of the link a person previously indicated between an attribute and that particular benefit in the AB-matrix. In fact, stating it in probabilistic terms, this specifically corresponds to the assumption that attributes and values are conditionally independent, given the benefits. If this assumption holds, it allows the researcher to collect information on AB- and BV-links in separate matrices, which is the cornerstone of the APT technique and its associated advantages. If this assumption of conditional independence is incorrect, the validity of APT as a technique to uncover means-end chains is in doubt.

In order to elucidate the issue of conditional (in)dependence, consider the example presented in Figure 3.2. It may represent part of a consumer's cognitive structure. Two ladders are distinguished, which share a common benefit ('healthy'). Whereas in the laddering interview all information is retained on the ladders 'lean beef' - 'healthy' - 'enjoying life' and 'hormone-free beef' - 'healthy' - 'safety', APT breaks the ladders in two parts: on the one hand the AB-links 'lean beef' - 'healthy' and 'hormone-free beef' - 'healthy', and on the other hand the two BV-links 'healthy' - 'enjoying life' and 'healthy' - 'safety' and these AB- and BV-links are treated as being independent of each other. The question however, is whether in reality, given a benefit, specific attributes have a high probability of leading to specific values, as shown in the left

panel of Figure 3.2. If this is not the case, as shown in the right panel of Figure 3.2, attributes and values are conditionally independent given the benefits.

Figure 3.2
Illustration of conditional dependence and independence of attributes and values for a given benefit



Means-end chain theory is not clear as to whether conditional independence is to be expected or not. This is probably due to the dominance of laddering as a data collection technique in means-end chain research. Laddering accommodates both conditional dependence and independence situations, so this assumption did not need to be considered before. However, Gutman (1982) already proposed separating AB- and BV-links. Walker and Olson (1991) also view the benefits as intermediating concepts, separating the product (the attributes) from the self (the values) without suggesting any dependence in the links, while Peter and Olson (1987) note that numerous combinations of attributes, benefit, and values are possible. In sum, we expect conditional independence to be supported.

3.3.2 Between-method Convergent Validity

Between-method convergent validity is an important aspect for any set of techniques, including methods to uncover means-end chains. In this

paper, we assess the convergent validity between APT and the 'standard' means-end chain measurement technique, laddering. Convergent validity is supported if cognitive structures identified with APT are similar to those obtained with laddering. We assess the similarity of APT and laddering with respect to the content as well as the structure of the means-end chains network.

The *content* of the means-end chains network refers to the nodes in the network. In this study, content is operationalized as the relative frequency at which the various attributes, benefits, and values occur. Complete similarity in content is attained if the concepts resulting from the laddering interviews and the responses in APT surface in equal proportions.²

We define the *structure* of the means-end chains network as the relative frequency of occurrences of the links (irrespective of the proportions in which the specific attributes, benefits, and values occur). Structure refers to the strength of the links between the concepts. If similarity in content is lacking, this does not necessarily imply a dissimilarity in structure, or vice versa. Thus, content and structure are two different aspects of convergent validity assessment.

If certain differences in content and structure of cognitive structures measured by APT and laddering arise, they may be the result of at least two sources: 1) the response mode chosen, and 2) the large and a priori fixed set of attributes in APT versus the small and variable set of attributes in which each laddering interview starts. Below we will elaborate on these two differences.

In at least two ways, the response mode chosen may cause differences between APT and laddering in the means-end chains identified. First, the fixed-format of APT reduces the natural flow of speech, which increases the extent of strategic processing by the respondent (Grunert and Grunert 1995). Strategic processes are those unwanted cognitive processes that do not result in activation of links

² Note that this definition also covers the situation where partially different sets of concepts occur in different methodologies, since then the frequency for a concept in one of the methods may equal zero.

between actual connected concepts, but occur when respondents choose alternative strategies just to arrive at an answer. On the other hand, the free-response format of laddering suffers from unwanted effects of interviewer involvement (Grunert and Grunert 1995). Second, the free-response format in laddering implies that respondents have to recall the concepts from memory, whereas the fixed-format of APT is associated with recognition of these concepts that are presented to the subjects. This may lead to differences in the content and structure of the means-end chains activated, as recall and recognition are quite different processes (e.g., Alba and Hutchinson 1987). The total number of recognized concepts typically exceeds the total number of recalled concepts (Singh and Rothschild 1983). This suggests that the total number of attributes, benefits, and values is greater in APT than in laddering. In principle this need not hamper the comparability of the results of APT and laddering, provided the impact does not vary over the set of attributes, benefits, and values or, more importantly, the links between them. If however, the impact differs, both the content and structure of the means-end chains identified with laddering and APT may diverge.

A second source of differences between laddering and APT data may be the set of attributes used as a basis for data collection. In APT, all respondents are confronted with the same set of attributes, which may consist of more than 10 items. In the laddering interview only a small number of attributes can be used (usually three to five), which can differ between respondents. Only these attributes and attributes elicited through (indirect) links with these attributes appear in the ladders. Since attributes are linked to benefits, which are again connected to certain values, the content of the more abstract levels may be influenced, while the structure may also be affected.

In sum, we identify at least two reasons why APT and laddering do not necessarily yield similar results. Clearly, for the progress in the area of means-end chains, it is important to know to what extent the two techniques yield comparable results. Additionally and equally important, it needs to be investigated whether APT is fundamentally valid, in that the attributes and values are conditionally independent, given the benefits. In the remainder of this chapter, we will empirically investigate

the two key validity issues of conditional independence and convergent validity.

3.4 Method

3.4.1 Data

Data are obtained from a laddering study on food products in Belgium. Three hundred laddering interviews are conducted, concerning yogurt (n=100), beef (n=100), olive oil (n=50) and vegetable oil (n=50). The interviews are equally divided over Flanders and Wallonia, which are the Dutch- and French-speaking parts of Belgium, respectively. All respondents bought the product in question and the laddering interviews were conducted by trained interviewers. In this laddering study we slightly deviated from the classical laddering procedure in that attribute elicitation involved different respondents than the laddering interview itself.

In order not to overlook any key characteristic of a product, a prespecified list of attributes was constructed by pretesting in a separate elicitation phase on 20 respondents for each product. Three approaches were jointly used to develop sets of salient product attributes: a thorough review of the literature, extensive questioning on the product, and a repertory grid task (Kelly 1955). The result of this final phase was a list of attributes for each product, which was used as point of departure in the depth-interview. Consumers were confronted with this list and were asked to indicate the importance of each attribute on a 3-point scale (not important to very important). This scale was used successfully in previous laddering research (Vanden Abeele 1992). Three very important attributes and two rather important attributes were chosen by the interviewer as starting point for the laddering procedure. The raw data were content analyzed and categorized, and subsequently the implication matrix was constructed.

In addition, APT-data for beef were collected. The attributes included in the AB-matrix consisted of the list of the twelve most

important attributes for beef used in the laddering study. The benefits and values included in the AB- and BV-matrices were selected on the basis of the frequency of their appearance in the implication matrix of the laddering study for beef. One hundred beef consumers completed the APT questionnaire, under the supervision of an interviewer. The AB- and BV-matrices included in the questionnaires are partially depicted in Figure 3.1.

3.4.2 Analysis

We tested the conditional independence assumption and the between-method convergent validity, using a loglinear model formulation. The loglinear model describes the probabilities that the concepts and links among them occur in ladders. Modeling the occurrence of ladders and links in a probabilistic framework is consistent with the ideas of Gutman (1991), who states that links “can be thought of in terms of some probability statement that one element will cause the occurrence of the other element (p.144).”

3.4.2.1 Conditional Independence

In order to test the assumption of conditional independence, we formulate a loglinear model representing the probability that a certain ladder occurs within a laddering interview. We start from the laddering data to test the assumption since it cannot by definition be tested on the APT data. From these data we form the three-way contingency table, indexed by A , B , and V , that represents all ladders. In Equation (3.1), the probability p_{ijk} that a ladder consists of attribute i , benefit j , and value k is expressed (on a logarithmic scale) as a function that is linear in the parameters. The set of parameters to be estimated consists of a scaling constant (α), the main effects, which capture the frequency of occurrence of attributes (β_i^A), benefits (β_j^B), and values (β_k^V), and their interactions ($\gamma_{ij}^{AB}, \gamma_{jk}^{BV}, \gamma_{ik}^{AV}$). The interactions capture the dependencies between attributes and benefits, benefits and values, and attributes and values, respectively, i.e.,

they represent the direct and indirect links between the concepts. The saturated model is formulated as follows:

$$(3.1) \quad \ln p_{ijk} = \alpha + \beta_i^A + \beta_j^B + \beta_k^V + \gamma_{ij}^{AB} + \gamma_{jk}^{BV} + \gamma_{ik}^{AV} + \delta_{ijk}^{ABV}$$

where δ_{ijk}^{ABV} is the error term. The product attributes (i) range from 1 to I , benefits (j) range from 1 to J , and values (k) range from 1 to K . In loglinear modeling terms, we write this model as $[ABV]$.³

Within this framework we can test for the conditional independence (Agresti 1990; DeSarbo and Hildebrand 1980) of attributes and values, given benefits. In APT, the links $A_i B_j$ and $B_j V_k$ are separated. Therefore, as discussed in the previous section, APT assumes that these links are independent. Thus, if we know that a specific benefit is linked to a particular attribute, this should not provide any information about the value to which this benefit leads. In a probabilistic framework this comes down to the assertion that the probability that a particular $B_j V_k$ -link is chosen (p_{jk}), does not depend on a preceding $A_i B_j$ -link. In terms of probabilities of links, conditional independence implies $p_{ijk} = (p_{ij} \cdot p_{jk}) / p_j$ (Agresti 1990), which corresponds to the loglinear model $[AB, BV]$. In order to test whether this model adequately represents the laddering data, we test it against the model that does specify the direct dependence of attributes and values in addition to their indirect dependence through the benefits. This latter model includes the interaction term AV : $[AB, BV, AV]$. The former model is nested in the latter model. If AV -dependence exists, it will appear in the laddering interview, since the interview does not impose independence, and information is retained on individual ladders. Therefore we use laddering data to test for the conditional independence of the attributes and values, given the benefits.

³ This means that whenever the loglinear model contains higher-order effects, it also incorporates lower-order effects composed from the variables. For example, the model $[AB, V]$ corresponds to the loglinear model with intercept, three main effects A , B , V , and the first order interaction between A and B .

The 'fit' of loglinear models can be evaluated by means of likelihood-ratio test statistics (χ^2), indicating the fit relative to the saturated model [ABV]. This statistic follows a chi-square distribution with degrees of freedom equal to the number of cells in the table minus the number of linearly independent parameters. In addition, model fit will be evaluated using Bonett and Bentler's (1983) normed fit index (Δ) which is independent of sample size. The index reflects the degree to which a model represents an improvement in goodness-of-fit (measured in χ^2) over the completely restricted model. When the model fully accounts for the data Δ takes on a value of 1.

Testing the conditional independence of the AB- and BV-links is similar to testing the hypothesis of no attribute-value interactions ($\gamma_{ik}^{AV} = 0$). The corresponding likelihood-ratio test statistic is the chi-square difference test statistic $\Delta\chi^2 = \chi_{AB,BV}^2 - \chi_{AB,BV,AV}^2$ (the difference in χ^2 between model [AB,BV] and model [AB,BV,AV]). The statistic is chi-square distributed with $(I-1)(K-1)$ degrees of freedom under the null hypothesis of conditional independence. If the test is not significant, the attributes and values are conditionally independent, which supports the procedure employed in APT of measuring AB- and BV-links in separate matrices.

3.4.2.2 *Between-method Convergent Validity*

The research question concerning the convergent validity between APT and laddering with respect to the content and structure of the means-end chains network, may also be formulated in terms of loglinear models. We collected both APT and laddering data for beef. We pool these data in order to test convergent validity, and form two 3-way contingency tables, one for the AB- and one for the BV-links. We introduce a new factor T in the loglinear models, representing the measurement technique from which the ladder originates: laddering ($t=1$) or APT ($t=2$). If the assumption of the conditional independence of attributes and values holds, we may treat the ladders as a sequence of independent links and formulate the (marginal) probabilities that links occur separately. In

loglinear modeling terms we define the probabilities of occurrence of links as:

$$(3.2) \quad \ln p_{ijt}^{ABT} = \alpha + \beta_i^A + \beta_j^B + \beta_t^T + \gamma_{ij}^{AB} + \gamma_{it}^{AT} + \gamma_{jt}^{BT} + \delta_{ijt}^{ABT}$$

$$(3.3) \quad \ln p_{jkt}^{BVT} = \alpha' + \beta_j^B + \beta_k^V + \beta_t^T + \gamma_{jk}^{BV} + \gamma_{jt}^{BT} + \gamma_{kt}^{VT} + \delta_{jkt}^{BVT}$$

Equations (3.2) and (3.3) are both saturated models.⁴ The probability that an AB- or BV- link appears in the data collected with measurement technique t is reflected by p_{ijt}^{ABT} and p_{jkt}^{BVT} , respectively. Whereas $\alpha, \alpha', \beta_i^A, \beta_j^B, \beta_t^T, \beta_k^V, \gamma_{ij}^{AB}$, and γ_{jk}^{BV} have the same interpretations as in Equation (3.1), β_i^T and β_t^T capture the difference in the overall frequency of occurrence of concepts between the two measurement techniques, i.e., attributes and benefits taken together in Equation (3.2), and benefits and values in Equation (3.3). The between-method difference in the frequency of occurrence of a specific attribute A_i , benefit B_j , or value V_k is taken into account by γ_{it}^{AT} , γ_{jt}^{BT} , γ_{jt}^{BT} and γ_{kt}^{VT} , respectively. In Equation (3.2) the between-method difference in the frequency of occurrence, or strength, of a specific $A_i B_j$ -link is captured by δ_{ijt}^{ABT} . Likewise, in Equation (3.3), δ_{jkt}^{BVT} reflects between-method differences for the $B_j V_k$ -links. According to our definition of content, the terms γ_{it}^{AT} , γ_{jt}^{BT} , γ_{jt}^{BT} , and γ_{kt}^{VT} represent the differences in content of the cognitive network, whereas δ_{ijt}^{ABT} and δ_{jkt}^{BVT} reflect the differences in structure.

The test of between-method similarity of the content of the means-end chains network may be expressed as $\{\gamma_{it}^{AT} = \gamma_{jt}^{BT} = 0\}$ for the AB-part, and $\{\gamma_{jt}^{BT} = \gamma_{kt}^{VT} = 0\}$ for the BV-part. This comes down to pairwise

⁴ A prime is added to the terms in Equation (3.3) that are also included in Equation (3.2) since the parameters need not have the same values.

testing of several restricted models. For the AB-part of the data (Equation (3.2)), we test for $\gamma_{it}^{AT} = 0$ and $\gamma_{jt}^{BT} = 0$ respectively, by testing the model $[A,B,T]$ against $[B,AT]$ and $[A,BT]$, all nested within Equation (3.2). Thus, we test the model which assumes no between-method differences in the relative frequencies of the various concepts versus the models which account for these differences. Additionally we may test for $\gamma_{it}^{AT} = 0$ and $\gamma_{jt}^{BT} = 0$ simultaneously by testing the model $[A,B,T]$ against $[AT,BT]$. Likewise, in Equation (3.3), these three tests come down to testing $[B,V,T]$ against the models $[V,BT]$, $[B,VT]$ and $[BT,VT]$. Analogous to the test of conditional independence, tests are based on the chi-square difference test statistic $\Delta\chi^2$, which is calculated from the difference in fit (χ^2) of each pair of models.

The similarity of structure may be assessed by testing $\{\delta_{ijt}^{ABT} = 0\}$ and $\{\delta_{jkt}^{BVT} = 0\}$. This corresponds to testing two restricted models in Equations (3.2) and (3.3), excluding the three-factor interaction terms δ_{ijt}^{ABT} and δ_{jkt}^{BVT} . Since the tests imply testing $[AB,AT,BT]$ and $[BV,BT,VT]$ against their saturated alternatives $[ABT]$ and $[BVT]$, the chi-square difference test statistic $\Delta\chi^2$ is identical to the χ^2 fit statistic of both models. Note here that by incorporating γ_{it}^{AT} , γ_{jt}^{BT} , γ_{jt}^{BT} and γ_{kt}^{VT} in Equations (3.2) and (3.3), we correct for the frequencies of the specific concepts in each measurement technique. If the chi-square difference statistic is significant for either the AB- or BV-part or both, we may conclude that APT and laddering do not reveal similar means-end chains in terms of structure.

3.4.2.3 Cell-sparseness

The likelihood-ratio test statistic used to assess conditional independence and convergent validity is sensitive to cell-sparseness (Margolin and Light 1974). Larntz (1978) indicated that the sample size should be at least five times the number of cells in the table. This condition may be violated in the context of laddering since many links may not occur in the

data, resulting in zero-entries in the contingency tables. Such zero-entries should be considered as sampling zeros and not as structural zeros, since there are typically no theoretical arguments supporting that such links cannot occur. Violation of the Larntz criterion causes the chi-square tests for nested models to be over-conservative. Therefore, in the cases where the Larntz criterion is not satisfied, we adjust the chi-square and chi-square difference test statistics. As suggested by Fienberg (1980) we use Williams' (1976) formula, which comes down to dividing the chi-square test statistics by a correction factor. See Williams (1976) for mathematical details.

3.5 Results

3.5.1 Conditional Independence

Since we are interested in testing the conditional independence of attributes and values, we focus on that part of the data where attributes and values come together at the intermediate level of the benefits. From the actual laddering interviews, we retain those ladders that contain complete attribute-benefit-value links, i.e. where an attribute is directly followed by a benefit and where this benefit is directly followed by a value. From the 100 interviews on beef and yogurt, and 50 interviews on olive oil and vegetable oil that were conducted, 128 ladders on beef, 148 ladders on yogurt, 60 ladders on olive oil and 47 ladders on vegetable oil, containing complete ABV-links, were retained. The number of different attributes, benefits, and values after content analysis and categorization are depicted in Table 3.1. The ladders are used to construct a 3-way $A_i B_j V_k$ contingency table for each product. The cells of the 3-way table contain the frequencies with which each of the links occurred in the data.

Table 3.1

Characteristics of the laddering data used to test for conditional independence

	Beef	Yogurt	Olive oil	Vegetable oil
# interviews	100	100	50	50
Attributes	10	10	8	5
Benefits	11	12	5	5
Values	13	17	8	5

Table 3.2

Tests for conditional independence of attributes and values given benefits

$[AB, BV, AV]$	Δ	Correction	χ^2_{adj}	df
Beef	.957	3.41	19.47	1080
Yogurt	.932	3.90	28.18	1584
Olive oil	.950	2.35	8.16	196
Vegetable oil	.962	1.77	4.35	64
$[AB, BV]$	Δ	Correction	χ^2_{adj}	df
Beef	.934	3.21	3.11	1188
Yogurt	.905	3.68	39.15	1728
Olive oil	.905	2.13	15.41	245
Vegetable oil	.919	1.64	9.42	80
$[AB, BV]$ vs. $[AB, BV, AV]$	Correction	$\Delta\chi^2_{adj}$	df	p-value
Beef	1.20	25.15	108	> .999
Yogurt	1.22	27.83	144	> .999
Olive oil	1.23	11.08	49	> .999
Vegetable oil	1.13	6.88	16	.975

Since in the four contingency tables the sample size is less than 5 times the number of cells, the condition of Larntz (1978) is violated. We therefore adjust the χ^2 test statistics according to Williams' formula (Williams 1976). Table 3.2 gives normed fit indices Δ , correction factors,

the adjusted χ^2 -statistics for both the unconditional $[AB, BV, AV]$ and conditional $[AB, BV]$ model, as well as the adjusted χ^2 -difference test-statistics for the tests of conditional independence.⁵ The adjusted χ^2 values for both the $[AB, BV, AV]$ and the $[AB, BV]$ models are low, indicating that the models fit the data very well. This is supported by the normed fit indices being close to one, indicating that there is not much improvement in fit possible. In each of the four tests, the assumption of conditional independence cannot be rejected. For all products, the p-value is very close to one. This provides strong empirical evidence for the independence of both the AB- and BV-matrices and therefore supports the partitioning of the ladders as in APT. Our findings indicate that attributes and values are associated indirectly through the attribute-benefit and benefit-value links.⁶

3.5.2 Convergent Validity

Since the independence assumption is supported, the models in Equations (3.2) and (3.3) can be estimated separately. In assessing the convergent validity of APT and laddering, we use all AB- and BV-links from the laddering interviews on beef and compare these data with links from the APT interviews on beef. We compare the number of entries in the AB- and BV-matrices in APT with the number of direct and indirect AB- and BV-links from the laddering implication matrix. For laddering, 462 links between attributes and benefits were observed and 1068 links between benefits and values. In APT these figures were 5773 and 9670, respectively.

⁵ Since the tables contain empty cells we have added a small constant (.5 as recommended by Goodman (1970)) to each of the cells. In this way the cells are treated as sampling zeros (as opposed to structural zeros).

⁶ The test of $[AB, BV]$ versus the saturated $[ABV]$ model is denoted as the unconditional test (Sobel 1995), and provides an alternative to the test of conditional independence $[AB, BV]$ versus $[AB, BV, AV]$ presented by us. For all products, the unconditional test rejects the hypothesis of conditional dependence, supporting the findings of the conditional test.

Both laddering and APT data were combined and organized in two 3-way contingency tables ($A_i B_j T_t$ and $B_j V_k T_t$) containing the frequencies of the $A_i B_j$ - and $B_j V_k$ -links, respectively. In both contingency tables, the ratio of observations to the number of cells exceeds 15. Hence, the condition suggested by Larntz (1978) is amply satisfied and the statistics do not need to be adjusted for cell sparseness as in Table 3.2. Table 3.3 contains the test statistics for the models nested within Equations (3.2) and (3.3). The statistical test results investigating the convergent validity of both approaches are given in Table 3.4.

Table 3.4 reveals that the tests for similarity in content with respect to the attributes ($[A,B,T]$ versus $[B,AT]$), benefits ($[A,B,T]$ versus $[A,BT]$), and attributes and benefits simultaneously ($[A,B,T]$ versus $[AT,BT]$) are all highly significant ($p < .001$). This suggests that the content of the AB-part of the means-end chains differs significantly between the two measurement techniques.

Table 3.3
Test statistics for nested model specifications

<i>Model</i>	Δ	χ^2	<i>Df</i>	<i>p-value</i>
[A,B,T]	.807	1994.69	337	<.001
[B,AT]	.821	1844.56	330	<.001
[A,BT]	.824	1810.03	315	<.001
[AT,BT]	.839	1659.91	308	<.001
[AB,AT,BT]	.982	180.49	154	.071
[B,V,T]	.726	3660.78	643	<.001
[V,BT]	.771	3063.83	624	<.001
[B,VT]	.748	3376.41	627	<.001
[BT,VT]	.792	2779.46	608	<.001
[BV,BT,VT]	.976	321.76	304	.232

Table 3.4

Tests for convergent validity of APT and laddering

Test	Test for	df	p-value	$\Delta\chi^2$
[A,B,T] vs. [B,AT]	Content	7	<.001	150.12
[A,B,T] vs. [A,BT]	Content	22	<.001	184.65
[A,B,T] vs. [AT,BT]	Content	29	<.001	334.78
[AB,AT,BT] vs. [ABT]	Structure	154	.071	180.49
[B,V,T] vs. [V,BT]	Content	19	<.001	596.95
[B,V,T] vs. [B,VT]	Content	16	<.001	284.37
[B,V,T] vs. [BT,VT]	Content	35	<.001	881.32
[BV,BT,VT] vs. [BVT]	Structure	304	.232	321.75

Similarity of structure is investigated by testing the model $[AB,AT,BT]$ against the saturated model. The $\Delta\chi^2$ -statistic is insignificant at the 5% level and the $[AB,AT,BT]$ model describes the data sufficiently well ($\Delta=.98$). Hence, for the AB-part of the data, we conclude that similarity in structure is supported.

Table 3.4 also shows the results of the model testing for the BV-part of the data. It shows the same picture as the AB-part. The three tests for the similarity of the content of the BV-part all reject the hypothesis at overwhelmingly low p-values. On the contrary, the $\Delta\chi^2$ -statistic is insignificant at the 5% level and the corresponding $[BV,BT,VT]$ model fits the data quite well ($\Delta=.98$). Thus, similarity in structure is supported for the BV-part.

We should note here that, statistically speaking, the lack of fit of models used to test the similarity in content ($[A,B,T]$, $[A,BT]$, $[B,AT]$, $[AT,BT]$, $[B,V,T]$, $[V,BT]$, $[B,VT]$ and $[BT,VT]$) is significant. Therefore the chi-square difference test statistics for similarity in content do not follow a central chi-square distribution, are formally inappropriate, and tend to reject the null-model too quickly. However, the extremely small p-values of the $\Delta\chi^2$ test statistics suggest that differences do indeed exist

in the content of the cognitive structures as measured by laddering and APT. The normed fit indices for these models are very close to the .8-level as suggested by Zahn and Fein (1979) as a cut-off level for models being acceptable. In addition, as discussed in section 3.3.2, there is strong empirical support for the differences in the frequencies with which the concepts occur.

3.6 Discussion and Conclusions

This chapter described APT as a quantitative technique for measuring means-end chains, and provided evidence on the validity of APT. Two major issues were addressed: the assumption of conditional independence of attributes and values given the benefits, and the convergent validity of APT as compared to the more traditional laddering interview.

The findings indicate that, across four product categories, attributes, and values are conditionally independent. This supports the APT approach in which AB- and BV-links are measured separately. This result is in line with the contention in means-end chain theory that the link between the benefits and the values concerns the self of the consumer and is independent of the product involved (Walker and Olson 1991).

With respect to between-method convergent validity, the results show that APT and laddering differ in the content (frequency of occurrence of specific concepts) of the means-end chains network identified. In general, APT yielded higher frequencies of occurrence of concepts than laddering. This is expected, given that APT involves recognition, and laddering recall. It cannot be stated a priori which technique yields the most valid information on content. In a recognition task, people may indicate more concepts than are really relevant to them. On the other hand, in a recall task, people may overlook important concepts, construct concepts during the interview, or be responsive to interviewer bias. Therefore, it is particularly important and reassuring that the strengths of the links (structure), corrected for the frequency of specific concepts, do not differ significantly between laddering and APT.

In summary, this chapter provides support for the validity of APT to uncover means-end chains. The findings do not suggest that laddering is not an important and useful technique for measuring means-end chains. Laddering was never intended to be used with large, representative samples. It is a qualitative technique with all its advantages and disadvantages. APT is geared towards being used in large-scale quantitative surveys in which representative samples are required (see for example chapter 4). However, applying APT requires exploratory research to identify the main concepts to be included in the AB- and BV-matrices. Laddering is the primary technique for this purpose.⁷ Using APT in conjunction with laddering is in spirit with other approaches in marketing research, such as compositional perceptual mapping and conjoint analysis, where qualitative techniques are used to identify key constructs that are subsequently summarized and used in fixed-format in large-scale studies (Green and Srinivasan 1978; Dillon, Frederick, and Tangpanichdee 1985).

The findings reported in this study indicate that APT is a viable quantitative technique for means-end chain research. This does not mean that APT is without its limitations. A criticism that may be levied against APT is the simplified representation of the means-end chains network by merely considering links between concepts at adjoining levels, i.e., the AB- and BV-links. This limitation may be alleviated by adding extra AA-, BB- and VV-matrices, containing the same concepts in the rows as in the columns. However, given similarity of the structures revealed with APT and laddering, the benefits of the additional information obtained with these matrices may be questioned, while in addition respondent burden increases substantially. Future research may shed more light on this issue.

Several other avenues for future research may be identified. First, another way to alleviate the potentially simplified A-B-V representation of APT is to extend the APT approach to the six-level hierarchy proposed

⁷ Alternatively, concept elicitation techniques like Kelly's repertory grid may be used to elicit both attributes and benefits. The values may be obtained from a value inventory like the List of Values (Kahle 1983).

by Peter and Olson (1987): concrete/abstract attributes, functional/psychosocial consequences or benefits, and instrumental/terminal values. This would increase the data collection burden for the respondent. Whether this is worth the effort is an issue for future research.

Second, in this study we analyzed the links between the attributes, benefits, and values ignoring whether ladders are elicited from the same or different subjects. A multi-level extension to the loglinear models proposed in this paper could be used to account for multiple ladders expressed by a single respondent. Recently, Poulsen and Valette-Florence (1996) proposed a heterogeneous latent Markov model to analyze means-end chains. The methodology enables the investigation of heterogeneity of means-end chains and, in that respect, extends the loglinear analysis performed in this study. Further research may examine the heterogeneity of perceived strengths of links between concepts and assess the validity of APT in a disaggregate setting.

Third, future research may focus on the underlying factors responsible for the differences in the occurrence of the concepts, especially the effects of recognition versus recall and the extent of unwanted strategic processes that may underlie APT. Finally, it is necessary to replicate this study with different products and in other countries to confirm and generalize our findings. In the next chapter, we will apply APT in an international segmentation study that seeks to link the product to the consumer at the cross-national segment level.

Chapter 4

International Market Segmentation Based on Consumer-Product Relations¹

4.1 Introduction

The globalization of the marketplace is arguably the most important challenge facing companies today (Yip 1995). Developments accelerating the trend toward global market unity include rapidly falling national boundaries, regional unification (e.g., European Union, NAFTA, ASEAN, Mercosur), standardization of manufacturing techniques, global investment and production strategies, expansion of world travel, rapid increase in education and literacy levels, growing urbanization among developing countries, free flow of information (e.g., World Wide Web), labor, money, and technology across borders, increased consumer

¹ This chapter is published as ter Hofstede, Frenkel, Jan-Benedict E.M. Steenkamp, and Michel Wedel (1999). "International Market Segmentation Based on Consumer-Product Relations," *Journal of Marketing Research*, 36 (February), 1-17.

sophistication and purchasing power, advances in telecommunication technologies, and the emergence of global media (Alden, Steenkamp, and Batra 1999; Hassan and Katsanis 1994; Mahajan and Muller 1994). Firms cannot ignore this trend toward globalization. Even if they decide not to be involved in the global (or pan-regional) marketplace, companies still face increased competition in their home markets due to nimble foreign competitors reaping the benefits of global strategies (Yip 1995). The marketing environment has become competitive to an extent that requires firms to target their products at markets spanning national boundaries. However, competitive clout can not be achieved in global consumer markets unless firms thoroughly understand and adequately respond to the core values and needs of those consumers (Hassan and Kaynak 1994).

Globalization affects consumer behavior and attitudes in a number of ways in that they transcend national borders. Groups of consumers in different countries often have more in common with one another than with other consumers in the same country. Therefore, there is a greater receptivity to global brands and foreign products across the world. Consumer brands have received wide global acceptance in categories such as consumer electronics, cars, fashion, home appliances, food products, and beverages (Hassan and Katsanis 1994). Many of these products respond to the needs and wants of segments of consumers that cut across national boundaries (Hassan and Kaynak 1994).

Thus, a major challenge facing today's international marketers is to identify global market segments and to reach them with products and marketing programs that meet the common needs of these consumers (Hassan and Katsanis 1994). This requires that characteristics of the (physical) product configuration are developed and marketed with the specific desires of the global target segment as a guiding principle. Such a consumer-oriented approach is critical to the success of global products and brands (Cooper 1984). Kleinschmidt and Cooper (1988) found that products designed for the global market achieve almost twice the market shares of products with domestic design, aimed at the same overseas markets. In a similar vein, Wind and Mahajan (1997) emphasize the importance of a global scope in product strategy. Research should be conducted in multiple countries rather than just in the home country (cf.,

Steenkamp, ter Hofstede, and Wedel 1999). This requires new developments in cross-national research methodology and models that integrate multi-country data at a pan-regional or global segment level (Wind and Mahajan 1997). It is in this area that the present study is situated.

This chapter proposes an integrated methodology to identify segments in international markets based on consumer means-end chains. As discussed in chapter 3, the key idea underlying means-end chains is that product attributes are *means* for consumers to obtain desired *ends*, i.e., values, through the benefits yielded by those attributes (Gutman 1982; Newell and Simon 1972, Reynolds and Olson 1998). In means-end chain (MEC) theory, these three concepts are hierarchically linked in cognitive structures in that product attributes yield particular benefits upon consumption, which contribute to value satisfaction. Means-end chain theory provides a conceptual basis for linking product and consumer in an international context, which facilitates successful product development and communication strategies (Gutman 1982). The interrelations among attributes, benefits, and values can differ across consumers within countries, but also across countries.

The methodology for identifying international consumer segments integrates data collection with a model for data analysis. The model is tailored to binary MEC data obtained with the Association Pattern Technique (APT; see chapter 3). Several features that are important for identifying valid international market segments are accommodated in our model.

1. The model captures the means-end relations that form the basis of cross-national segments in a stochastic framework. In the model probabilistic relations between attributes, benefits, and values are specified at the segment level.
2. The model captures differences in response behavior that may exist between consumers in different countries, which is an important issue in comparing and grouping data across countries (Chen, Lee, and Stevenson 1995; Steenkamp and Baumgartner 1998). The model also captures within-country heterogeneity in response behavior that may exist among consumers.

3. In modeling the segment sizes, we use a specification where country memberships enter as concomitant variables. This enables the segment sizes to differ across countries, so that truly pan-regional or global segments are accommodated.
4. We use a pseudo maximum likelihood approach proposed in chapter 2, to arrive at consistent estimates of the model parameters. Since international marketing research generally deals with samples stratified by nation with equal allocation of sample sizes, traditional estimation models result in biased parameter estimates.

Whereas the model is specifically suitable for international market segmentation, it nests situations of domestic segmentation, i.e., segmentation in a single country. In the discussion of this chapter, we show how domestic segmentation is accommodated as a special case of our model.

To assess the performance of the model, a Monte Carlo analysis of synthetic data is included. The model is applied to data from a large-scale international consumer survey sponsored by the European Commission. The sample comprises nearly 3,000 consumers from eleven countries of the European Union. The results and their implications for marketing strategy are discussed. The segments are related to descriptor variables and the predictive validity is assessed. The incremental contribution of the model over traditional approaches is evaluated both conceptually and empirically. Finally, conclusions, limitations, and issues for further research are presented.

4.2 The International Segmentation Basis

In international segmentation studies, information at the country level often has been used as basis for grouping countries into geographic segments. Macro-level geographic, political, economic, and cultural data typically have been used to identify market segments consisting of groupings of countries (e.g., Helsen, Jedidi, and DeSarbo 1993; Kale 1995). This approach provides insights into which groups of countries potentially can be targeted, but it provides no information on which

consumers in those countries will respond to marketing effort. Within-country heterogeneity is ignored, in spite of being the very basis for identification of domestic segments. If consumers instead of countries were used as the basis for identifying international segments, the effectiveness of marketing strategies would increase greatly (Jain 1989).

Recently, at least two broad classes of micro-level segmentation bases have been used in international segmentation research. Some studies have used product-specific bases including perceived product characteristics (e.g., Moskovitz and Rabino 1994; Yavas, Verhage, and Green 1992). Other studies have utilized consumer-specific bases such as lifestyles (e.g., Boote 1983) and values (e.g., Kamakura et al. 1993). Whereas consumer-specific bases are closer to the consumer, product-specific bases are more actionable (Shocker and Srinivasan 1979). Building upon those studies, we propose to incorporate both consumer and product characteristics in international market segmentation. Using measures that connect the consumer and the product creates the potential to increase the effectiveness of global or pan-regional marketing strategies. Such an approach helps to link product development with communication strategies that attempt to position products by associating the product aspects to the achievement of desired ends (Gutman 1982).

Newell and Simon (1972) developed the concept of the means-end chain (MEC). MEC explicitly establishes the relation between product and consumer by positing hierarchical links between product attributes, the beneficial aspects of product use, and consumers' values: attributes lead to benefits which contribute to value satisfaction. In MEC theory, the dominant meaning of a product attribute to a consumer is determined by the benefits it is perceived to lead to. Benefits accrue to a consumer when the product is purchased and used (Peter and Olson 1993). They derive valence and importance from their perceived ability to satisfy personal values. Values underlie a large and important part of human cognition and are among the most central determinants of consumer behavior (Pitts and Woodside 1984; Steenkamp, ter Hofstede, and Wedel 1999; see also chapters 2 and 3). They transcend specific objects, in contrast to attributes and benefits, which relate to a particular product (Rokeach 1973). In sum, MEC theory posits that the way in which physical product attributes

are linked to personal values defines how products gain personal relevance and meaning. An attribute is important if it leads to a desired benefit, while the perceived benefit derives its importance through the extent to which it is linked to one or more personal values (Reynolds and Olson 1998).

Linking products and consumers in means-end chains for the purpose of international segmentation combines the strengths of product-specific and consumer-specific bases of segmentation. Identifying segments from concrete product attributes increases the actionability of the results for product development (Urban and Hauser 1993). On the other hand, it is well known that consumers buy products for what these products can do for them and how they contribute to achieving desired ends, not for their physical attributes per se. Thus, effective marketing strategies require insights into the links between product attributes, benefits, and values. The benefits associated with product attributes support persuasive claims of the product. Related values help to substantiate and frame these claims in targeted advertising campaigns (Young and Feigin 1975). Hence, it is not surprising that a number of researchers have suggested that means-end chains provide a particularly suitable basis for market segmentation (Gutman 1982; Kamakura and Novak 1992).

4.3 The International Segmentation Methodology

4.3.1 Model Formulation

In our approach, the perceived means-end relations between attributes, benefits, and values are measured at the individual level using the Association Pattern Technique (APT) described in chapter 3. The underlying idea of APT is to decompose a MEC relation into two independent parts, comprising attribute-benefit (AB) links and benefit-value (BV) links, respectively. The APT data collection procedure is extensively validated in chapter 3.

In modeling the data collected in an APT-task, we consider the consumer's response process in a probabilistic framework. This is consistent with the ideas of Gutman (1991), who states that links "can be thought of in terms of some probability statement that one element will cause the occurrence of the other element (p. 144)." The choice of a link in either of the two APT matrices depends on the intrinsic strength of the link. It also depends on the respondent's response tendency, i.e., his/her propensity to indicate a link in general, regardless of the underlying strength (importance) of the AB or BV relation. Measurement instruments are subject to response tendencies and those tendencies may differ across countries (Chen, Lee, and Stevenson 1995; Steenkamp and Baumgartner 1998). International segmentation studies employing consumer surveys are affected adversely by cross-national differences in response behavior. Variation in response tendencies may mask cross-national segments, and the segments identified could constitute groups of consumers that differ in response tendency rather than in the underlying behavior studied. In this study, we explicitly account for differences in response tendency within and across countries.

In our model segments are represented by a finite mixture formulation (Titterton, Smith, and Makov 1985). Let u_j^m denote the intrinsic strength of link j in matrix m ($j = 1, \dots, J_m$), and y_{ij}^m a binary variable with $y_{ij}^m = 1$ if subject i ($i = 1, \dots, I$) chooses link j , in the AB ($m = 1$) or the BV ($m = 2$) matrices and $y_{ij}^m = 0$ otherwise. Analogous to item response theory (Baker 1992), it is assumed that if the strength of link j exceeds subject's i 's threshold ($u_j^m > \theta_i$), s/he chooses the link. Thus, the smaller the threshold, θ_i , the higher the propensity of the subject to choose any link, whereby θ_i captures the response tendency of subject i . The probability of subject i choosing link j can then be expressed as a function of the strength u_j^m and the threshold parameter θ_i . We assume the cumulative distribution of u_j^m to be logistic, so that we

obtain a closed-form expression for the probability of observing a link j for subject i :

$$(4.1) \quad P[y_{ij}^m = 1 | u_j^m, \theta_i] = P[u_j^m > \theta_i] = \frac{\exp[u_j^m - \theta_i]}{1 + \exp[u_j^m - \theta_i]}.$$

For reasons of identification, the threshold parameters are restricted to sum to zero, i.e., $\sum_i \theta_i = 0$. Equation (4.1) corresponds to the Rasch (1960) model as used in item response theory (Baker 1992). Equation (4.1) shares the following properties with the Rasch model. Both u_j^m and θ_i are defined on the real axis. A higher u_j^m means that the particular link is stronger, and the higher θ_i , the lower the propensity that a person chooses links in general. The distance between u_j^m and θ_i may be interpreted in terms of log-odds ratios, i.e., the logarithm of the ratio of the probability of observing a particular link to the probability of not observing that link.

We use the strength of the link as the basis for segmentation. For that purpose we let u_j^m vary over S unobserved segments ($s = 1, \dots, S$). Given segment s , the subject-specific density is:

$$(4.2) \quad f_{i|s}(y_i | \theta_i; u_s) = \frac{\exp \left[\sum_{m=1}^2 \sum_{j=1}^{J_m} y_{ij}^m (u_{js}^m - \theta_i) \right]}{\prod_{m=1}^2 \prod_{j=1}^{J_m} [1 + \exp(u_{js}^m - \theta_i)]},$$

where $u_s = [u_{js}^m]$ and $y_i = [y_{ij}^m]$.

There are, however, several problems associated with this formulation. The fixed-effects estimator of the incidental threshold parameters θ_i is inconsistent (Ghosh 1995) and not identified at the individual level (Lindsay, Clogg, and Grego 1991). Yet, not accounting for the heterogeneity in response tendencies may confound the estimates of the segment-specific link-strengths with response tendencies. A possible solution would be to let θ_i vary across the S unobserved segments, i.e., $\theta_i = \theta_s$. However, this approach has several disadvantages.

It is not evident a-priori that the thresholds should obey a similar segment structure as the link parameters. In addition, for many applications such an approach would be too restrictive in the sense that it does not deal with within-segment heterogeneity since all subjects within a segment are constrained to have an identical threshold. Research has shown that the assumption of within-segment homogeneity may be overly restrictive, resulting in a loss of explanatory and predictive power (Allenby, Arora, and Ginter 1998). Most important, the segments identified through such a model would differ in their response tendencies. It is exactly this confounding of response behavior and the structural basis of the segments (i.e., the strengths of the links) that we want to avoid. Therefore, we treat the thresholds as random parameters that vary across the population according to a normal distribution. We employ a specification of the thresholds that separates within-country and between-country heterogeneity, by assuming that the thresholds θ for country c ($c = 1, \dots, C$) have normal distributions, $\Phi(\theta; \mu_c, \nu_c)$, with mean μ_c and variance ν_c . Whereas Equation (4.2) presents the conditional distribution of subject responses, given θ , the unconditional distributions are obtained by taking the expectation of Equation (4.2) over the thresholds:

$$(4.3) \quad f_{i|s}(y_i; u_s, \mu_{c(i)}, \nu_{c(i)}) = \int_{-\infty}^{+\infty} f_{i|s}(y_i | \theta; u_s) d\Phi(\theta; \mu_{c(i)}, \nu_{c(i)}),$$

where $c(i)$ indicates the country membership of subject i . The model simultaneously accounts for between-country differences in response thresholds and within-country threshold heterogeneity.

The expression in Equation (4.3) is conditional upon segment s . Letting $\pi_{s|c}$ denote the proportion of consumers in country c belonging to segment s , and using a standard finite mixture formulation, the unconditional distribution (on s) equals

$$(4.4) \quad f_i(y_i; u, \mu_{c(i)}, \nu_{c(i)}) = \sum_{s=1}^S \pi_{s|c(i)} f_{i|s}(y_i; u_s, \mu_{c(i)}, \nu_{c(i)}).$$

Since segment sizes may differ between countries, i.e., a particular segment is more prevalent in some countries than in others, the segment

proportions $\pi_{s|c}$ are allowed to vary across countries through a concomitant variable specification (Wedel and Kamakura 1998, p. 147):

$$(4.5) \quad \pi_{s|c} = \frac{\exp(\gamma_{sc})}{\sum_{t=1}^S \exp(\gamma_{tc})}.$$

By imposing the restrictions $\gamma_{sc} = 0$ within each country, the proportions sum to one across segments.

4.3.2 Estimation and Model Comparison

As discussed in chapter 3, international samples are often stratified by country, with approximately the same sample size in each country. This results in sample sizes that are not proportional to population sizes. Since conventional maximum likelihood estimates may be severely biased when applied to stratified samples, we use the pseudo-maximum likelihood (PML) approach proposed in chapter 2 (see also Skinner, Holt, and Smith 1989). The unconditional distribution presented in Equation (4.4) equals the likelihood contribution of subject i , L_i . The sample log-pseudo likelihood is defined as:

$$(4.6) \quad \ln L_P = \sum_{i=1}^I w_i \ln L_i,$$

with the sampling weights w_i being equal to the ratio of the population share to the sample share as in section 2.2.2 on stratified samples. Maximizing the log-pseudo likelihood in Equation (4.6) with respect to the unknown parameters γ_{sc} , u_{js}^m , μ_c , and v_c , yields consistent estimates of the population parameters. We use a quasi-Newton algorithm to maximize the log-pseudo likelihood in Equation (4.6). Gauss-Hermite integration of order $R = 32$ is used to evaluate the integral in Equation (4.3). In that way, Equation (4.3) reduces to a weighted sum, so that the score functions may be derived in closed form (see Appendix 4.B). A problem related to the numerical maximization of the likelihood is the

occurrence of local optima. The solution chosen here is to run the models from a range of randomly chosen starting values for the parameters.

To determine the appropriate number of segments, likelihood-ratio test statistics cannot be used. The LR-statistic is not asymptotically distributed as chi-square since certain regularity conditions do not hold (Aitkin and Rubin 1985). Instead, we use the pseudo information criteria proposed in chapter 2 for selecting the appropriate number of segments. We use the Pseudo Bayesian Information Criterion, PBIC, and the Pseudo Consistent Akaike Information Criterion, PCAIC. They are defined by $PBIC = -2\ln L_P + K\ln I$ and $CAIC = -2\ln L_P + K\{\ln(I)+1\}$, where I denotes the number of observations and K the number of free parameters in the model ($K = S(J+C)+C-1$). We choose that number of segments for which these modified criteria have a minimum value.

These criteria are also used to test for between and within-country heterogeneity in response behavior. In testing for the existence of between-country heterogeneity, the pseudo-information criteria, PCAIC and PBIC, of the (unrestricted) model in Equation (4.6) are compared to those of a restricted model where $\mu_1 = \mu_2 = \dots = \mu_C = 0$. Since the overall mean of the thresholds has already been restricted to zero, the number of free parameters under this restricted model is $S(J+C)$. In testing for the existence of within-country heterogeneity, the unrestricted model is compared to the model with threshold variance restricted to zero ($v_1 = v_2 = \dots = v_C = 0$), which has $S(J+C) - 1$ free parameters.

4.3.3 Country-specific Entropy

In order to assess how well each country fits the international segmentation scheme, we propose to use a normed entropy measure (Ramaswamy et al. 1993). We define the entropy measure for each country separately: $E_{s|c} = 1 + \sum_s \sum_{i:c(i)=c} \alpha_{is} \ln \alpha_{is} / (I_c \ln S)$, where I_c is the number of subjects in country c , and α_{is} the posterior probabilities obtained by Bayes' rule:

$$(4.7) \quad \alpha_{is} = \frac{\pi_{s|c(i)} f_{i|s}(y_i; u_s, \mu_{c(i)}, v_{c(i)})}{\sum_{t=1}^S \pi_{t|c(i)} f_{i|t}(y_i; u_t, \mu_{c(i)}, v_{c(i)})}.$$

The entropy for a particular country approximates one if the posteriors for all consumers in that country are close to either zero or one. If the α_{is} in country c are all near $1/S$, $E_{S|c}$ tends to zero. A country with lower entropy has a higher proportion of subjects that are difficult to classify in one of the international segments. In this case, a substantial number of subjects within that country have specific means-end chain structures that do not agree with those of one of the segments. Hence, deviations from the pan-regional or global segmentation scheme may be identified. Note that the unconditional entropy measure E_S is biased in the case of stratified samples, i.e., large (small) countries will have a too low (high) contribution to the entropy. The country-specific entropy is justified, because of its independence of the stratified sample design.

4.4 Monte Carlo Analysis

To assess the performance of the model, we performed a Monte Carlo study in which we applied the model to synthetic data with known parameter values. The design of the Monte Carlo study is based on the following four factors that are hypothesized to affect parameter recovery:

PRB: The strength of link-probabilities, $p_{js}^m = \exp(u_{js}^m) / [1 + \exp(u_{js}^m)]$,

which are generated either in the .3-.7 or the .1-.9 range;

TSE: The within-country dispersion of thresholds, with standard errors

$\sqrt{v_c}$ generated in the range of .3-.5 or .5-.8, respectively;

SEG: Number of segments, $S=2$ or $S=5$.

DIS: The distribution underlying the strength of the links, logistic and normal, which represent situations where the distribution of the model is correctly specified versus misspecified.

We used a full 2^4 factorial design to generate data for 9 countries, with 300 respondents per country, so that we had 2700 observations for each model run. In each cell of the design 10 replications were generated, resulting in 160 data sets to be analyzed. The total number of links in the APT matrices was 18.

To generate the link probabilities, the u_{js}^m were taken from a fixed interval in such a way that the p_{js}^m were confined to the range dictated by the particular experimental condition. Then, according to the level of *DIS*, logistic or normal distributed errors were added to the u_{js}^m . The thresholds were drawn from a normal distribution, with standard errors $\sqrt{v_c}$ for each country, in the range dictated by the particular experimental condition. The μ_c were taken from an interval such that the overall probabilities (given both u_{js}^m and μ_c) of choosing a link were in the interval [.05, .90]. The prior segment proportions $\pi_{s|c}$ were taken from the interval [.15, .60]. Then, the y_{ij}^m were generated as binary variables equaling one when $u_{js}^m > \theta_i$ and zero otherwise (for all i and j , and i in segment s).

Our model is used to analyze each of the 160 synthetic data sets, using 10 random starts to overcome local optima. As measures of model performance we calculated the root mean squared error (RMSE) of the estimated versus the true parameters for: p_{js}^m , $\pi_{s|c}$, μ_c , and $\sqrt{v_c}$, denoted by $RMSE(p)$, $RMSE(\pi)$, $RMSE(\mu)$, and $RMSE(\sqrt{v_c})$, respectively. In addition, to assess how well the model identifies the segments from the data, we calculated the percentage of subjects that were correctly assigned to the segments (denoted by *PRED*). The optimal-Bayes rule of classification is used, which assigns each subject to the segment with the largest posterior segment membership probability. The log or logit transformed RMSE's (and predictions) were subjected to analyses of variance to investigate the main effects and the interactions of *PRB*, *TSE*, *SEG*, and *DIS*.

For all measures, the F-ratios of the main and interaction effects, their significance and R^2 -values are reported in Table 4.1. We report the marginal means for each level of the factors in Table 4.2. The R^2 -values in Table 4.1, show that, as compared to the other measures of model performance, the variation of *PRED* is very well explained by the experimental design ($R^2 = .991$), whereas the variation explained in

RMSE(μ) is relatively low ($R^2 = .273$). The R^2 -values for the other measures are close to .6.

Table 4.1

ANOVA results for model performance measures

<i>Factor</i>	<i>F-ratio</i> <i>PRED^a</i>	<i>F-ratio</i> <i>RMSE(p)^a</i>	<i>F-ratio</i> <i>RMSE(π)^a</i>	<i>F-ratio</i> <i>RMSE(μ)^b</i>	<i>F-ratio</i> <i>RMSE</i> <i>($\sqrt{v_c}$)^b</i>
<i>PRB</i>	7556.35*	34.17*	103.33*	7.58*	51.05*
<i>TSE</i>	5.95	2.16	.39	1.23	33.92*
<i>SEG</i>	7295.32*	174.23*	98.29*	.05	2.36
<i>DIS</i>	42.48*	6.73	1.47	36.31*	87.03*
<i>PRB</i> $\times$ <i>TSE</i>	3.04	.30	1.72	.76	.27
<i>PRB</i> $\times$ <i>SEG</i>	884.13*	9.79*	.05	.06	.27
<i>PRB</i> $\times$ <i>DIS</i>	18.31*	.22	4.08	.81	.24
<i>TSE</i> $\times$ <i>SEG</i>	1.29	.05	.07	.06	.98
<i>TSE</i> $\times$ <i>DIS</i>	.42	.83	.60	.25	.11
<i>SEG</i> $\times$ <i>DIS</i>	7.10*	1.34	1.48	1.73	2.08
<i>PRB</i> $\times$ <i>TSE</i> $\times$ <i>SEG</i>	1.14	.33	.01	.25	.68
<i>PRB</i> $\times$ <i>TSE</i> $\times$ <i>DIS</i>	.47	1.57	.88	.03	.13
<i>PRB</i> $\times$ <i>SEG</i> $\times$ <i>DIS</i>	6.52	.66	3.09	.06	1.17
<i>TSE</i> $\times$ <i>SEG</i> $\times$ <i>DIS</i>	.00	.23	.08	.15	.90
<i>PRB</i> $\times$ <i>TSE</i> $\times$ <i>SEG</i> $\times$ <i>DIS</i>	1.23	.39	4.12	4.72	.17
R^2	.99*	.62*	.60*	.27*	.56*

^a transformed by logit

^b transformed by log

* $p < .01$

Recovery of segment assignment (PRED) is significantly affected by the link probability interval (PRB), the number of segments (SEG), and the distribution of the strength of the links (DIS). The measure improves with increasing link probability intervals and decreases with the number of

segments (Table 4.2), which is to be expected. However, the correct assignment in the S=5 condition of 64.2% of the cases is still high, as compared to a random assignment of 20%. Recovery of segment assignment decreases only slightly when the distribution is misspecified (a decrease from 77.8% to 76.4%). The PRB×SEG interaction in the ANOVA of PRED is the only sizable significant interaction in Table 4.1. Inspection of the marginal means in the PRB×SEG table indicates that the improvement of the predictions due to more variation in the link probabilities is much higher in the case of the S=5 solution, as compared to the S=2 solution. Apparently, when the number of segments increases, more distinctive patterns of the link probabilities improve classification of subjects to the segments more.

Table 4.2

Mean model performance measures according to design factors

<i>Factor</i>	<i>Level</i>	<i>PRED</i> (%)	<i>RMSE</i> (<i>p</i>)	<i>RMSE</i> (π)	<i>RMSE</i> (μ)	<i>RMSE</i> ($\sqrt{v}$)
<i>PRB</i>	[.3, .7]	63.9*	.070*	.070*	.186*	.131*
	[.1, .9]	9.4	.050	.040	.158	.098
<i>TSE</i>	[.3, .5]	77.3	.056	.054	.166	.129*
	[.5, .8]	76.9	.063	.056	.178	.100
<i>SEG</i>	S=2	9.0*	.031*	.039*	.171	.111
	S=5	64.2	.089	.071	.173	.118
<i>DIS</i>	logistic	77.8*	.057	.053	.140*	.093*
	normal	76.4	.062	.057	.204	.136

* Differences between means significant on log or logit scale ($p < .01$)

Analogous to PRED, recovery of the link probabilities, $RMSE(p)$, and the segment proportions, $RMSE(\pi)$, are significantly affected by PRB and SEG. They improve with larger intervals of the link probabilities and with fewer segments. Apparently, an increase in the probability interval and less segments helps the identification of the segments, and therefore,

improves the recovery of these parameters. The magnitude of the threshold variance (TSE) and the distributional assumption (DIS) have no significant influence on $RMSE(p)$ and $RMSE(\pi)$, which supports the robustness of the approach.

The recovery of the country-specific mean of the threshold parameters, $RMSE(\mu)$, is significantly affected by PRB and DIS. Its recovery improves with larger intervals of the link probabilities (PRB) and correct distributional assumptions (DIS). The estimates of the standard errors of the thresholds, $RMSE(\sqrt{v_c})$, are affected by PRB, TSE and DIS. Larger intervals of the link probabilities (PRB), larger variances of the thresholds (TSE), and correct underlying distribution of the link strengths (DIS) have a positive significant effect on the recovery of the threshold variances. The precision of the threshold parameters, $RMSE(\mu)$ and $RMSE(\sqrt{v_c})$, are not affected by the number of segments. This is caused by the number of threshold parameters being independent of the number of segments, as opposed to $RMSE(p)$ and $RMSE(\pi)$.

Most important, the overall means of the measures $RMSE(p)$, $RMSE(\pi)$, $RMSE(\mu)$, $RMSE(\sqrt{v_c})$ are low: .06, .05, .17, and .11, respectively, and the mean of $PRED$ is high (77%). This means that the model recovers the true parameter values well and is effective in correctly assigning subjects to the segments. In general we find that the structural part of the model is robust against misspecification of the distribution and response variance. As might be expected, the model performs better when the link probabilities have larger intervals, in particular for large numbers of segments. In our experience with APT in empirical studies, we typically find large variation in the link probabilities.

4.5 Empirical Application

4.5.1 Data Collection

Data were collected in an international survey in the EU, as part of a major study on foods conducted for the European Commission. For the purpose of the application of our model, APT data on yogurt are used.

The dairy industry is among the biggest food industries in the EU. On average, European consumers spend close to 15 percent of the overall food budget on dairy products and annual sales of yogurt in the EU are about \$10 billion. The yogurt market has been revitalized in the 1990s through new product development activity and heavy investments in branding and market communication, among others. This has led to increasing differentiation of the market. The best market opportunities are expected from luxury/indulgence yogurts and enriched yogurts, while plain yogurt is declining. Yogurt is marketed by large multinational companies such as Nestle, Danone, and Sodialal-Yoplait, which operate at a European scale, as well as by smaller national companies. Companies have also set up joint ventures and strategic alliances to increase their market power in the yogurt market. See Valli, Traill, and Loader (1997) for more details.

For data collection, mail questionnaires were sent out to households in eleven countries of the European Union: Belgium, Denmark, France, Germany, Great Britain, Greece, Ireland, Italy, the Netherlands, Portugal, and Spain, i.e., all EU-12 countries except for Luxembourg.² The fieldwork was carried out in 1996 by a pan-European market research agency. The questionnaires included the APT-task for yogurt, consumption patterns, socio-demographic, personality and attitudinal variables as well as additional measures that are not of interest to the purposes of this study.

Before collecting the pan-European data, extensive cross-national pretests were conducted. First, as suggested in chapter 3, the relevant attributes and benefits were elicited by 100 in-depth laddering interviews on consumers' MEC-structures for yogurt. The resulting data were content analyzed into attributes, benefits, and values. These attributes, benefits, and values were included in APT matrices that were administered to consumers in Denmark, Greece, and Great Britain (N=150). Based on frequency of occurrence, eight attributes and ten benefits were selected and included in the APT matrices that were used in

² The EU started in 1992 with 12 countries, which are sometimes denoted as the EU-12 countries. Since then Sweden, Austria, and Finland have joined the EU.

the main study. The values used in APT are those from the 'List of Values' (LOV) value inventory (Kahle 1986). The LOV contains nine personal values that are relevant for consumer behavior (Kahle 1986; Kamakura and Novak 1992) and was used in international context in the empirical application of chapter 2. These APT matrices were pre-tested among 48 Dutch consumers for wording, interpretability, and layout. Appropriate adjustments were made. Subsequently, an additional pretest was conducted in an international context using samples in France, Great Britain, Greece, the Netherlands, and Spain ($N = 164$). Again, appropriate adjustments were made and the final concepts included in the APT matrices were selected. Those concepts are shown in the APT matrices given in Appendix 4.A. In the whole process, back-translation methods were used to ensure a similar content of the statements in the languages involved.

A sample, stratified by country, was drawn randomly from the household consumer panels of the pan-European market research agency. These panels are representative of the national populations with respect to a large number of socio-demographic characteristics. After sending reminders, the overall response was around 70 percent; 3,147 questionnaires were obtained from which 2,961 were usable for analysis (see Table 4.3).

Table 4.3
Sample characteristics

	<i>Response</i>	<i>Sample size^a</i>	<i>Population size^b</i>	<i>Weights w_c</i>
Belgium	70%	272	4,000	.343
Denmark	82%	229	2,374	.242
France	82%	346	21,700	1.462
Germany	80%	261	32,200	2.876
Great-Britain	69%	254	22,675	2.081
Greece	75%	292	3,204	.256
Ireland	74%	211	1,100	.122
Italy	59%	276	19,000	1.605
Netherlands	76%	260	6,504	.583
Portugal	45%	267	3,060	.267
Spain	69%	293	11,221	.893

^a Number of households

^b × 1000 households

4.5.2 Results

We applied the segmentation model to the APT yogurt data. The model is estimated for $S=2$ to $S=5$. To overcome local optima, we ran each model from ten sets of random starting values. The solution with maximum log-pseudo likelihood is reported. The information criteria for segment-solutions $S=2$ to $S=5$ are given in Table 4.4. Both PBIC and PCAIC have minimal values for $S=4$, supporting the fact that the four-segment solution provides the most parsimonious representation of the data. We denote the four segments as $S1$, $S2$, $S3$, and $S4$. The entropy measures $E_{S|c}$ are shown in Table 4.5. All values exceed .8, which indicates that in none of the countries substantive segments have been overlooked, and that the segments are well separated.

Table 4.4
Segment selection criteria

<i>S</i>	<i>Log-pseudo likelihood</i>	<i>PCAIC</i>	<i>PBIC</i>
2	-167,908	338,657	338,341
3	-166,537	337,291	336,822
4	-164,950	335,494*	334,872*
5	-164,708	336,385	335,610

* Denotes minimum value

In Table 4.5, we present the estimates of the country-specific segment proportions $\pi_{s|c}$. The sizes of segments *S1*, *S2*, and *S3* range from 15 to 20 percent. Segment *S4* is by far the largest segment, comprising almost half of the population. Segment *S4* is represented in substantial proportions in all countries in the sample and is therefore a truly pan-European segment. Segments *S1*, *S2*, and *S3* are cross-national segments, being present in substantial proportions in multiple countries. These segments are not truly pan-European since they are not substantially represented in all countries. Segment *S1* is mainly located along the coastal periphery of Europe (Denmark, Great Britain, Ireland, the Netherlands, Portugal, Greece, and Spain), consumers from Germany mainly dominate segment *S2*, and *S3* is mainly located in the north-west of Europe (Great Britain, Germany, Denmark, Ireland, and the Netherlands).

The estimates of the response parameters are shown in the last two columns of Table 4.5. The mean-thresholds, μ_c , vary over countries, being the highest in Great Britain and Denmark and the lowest in Italy and France. This means that, on average, in Great Britain and in Denmark respondents are least inclined to indicate links in the APT task, and in Italy and France they are most inclined to do so. The within-country threshold variances, v_c , are similar in magnitude across countries. They have the smallest value for the Netherlands (.318) and the largest value for Denmark (.487); the Danish appear somewhat less homogeneous in their response tendency and the Dutch somewhat more. As compared to the absolute values of the mean thresholds, μ_c , the magnitude of the v_c 's

is large, indicating that within-country heterogeneity in response thresholds is larger than between-country heterogeneity.

Table 4.5

Segment sizes and threshold estimates

	$E_{s c}$	$\hat{\pi}_{1 c}$	$\hat{\pi}_{2 c}$	$\hat{\pi}_{3 c}$	$\hat{\pi}_{4 c}$	$\hat{\mu}_c$	$\hat{v}_c$
Belgium	.876	12.2	17.5	8.5	61.8	-.017	.350
Denmark	.846	47.6	2.6	27.0	22.8	.173	.487
France	.899	2.4	21.8	3.3	72.4	-.145	.320
Germany	.826	3.2	45.1	26.3	25.4	-.018	.388
Great Britain	.807	41.2	7.0	26.0	25.9	.286	.339
Greece	.823	28.2	13.4	6.5	52.0	-.056	.466
Ireland	.806	40.6	12.1	17.9	29.5	-.045	.411
Italy	.896	5.9	10.2	5.1	78.8	-.228	.480
Netherlands	.821	31.5	14.3	17.9	36.3	.040	.318
Portugal	.835	35.2	27.8	4.9	32.1	-.009	.338
Spain	.832	25.1	8.5	3.8	62.6	.019	.361
Total*		17.4	20.9	15.8	46.0		

* The total is calculated as a weighted average of the country-specific priors.

The information criteria PCAIC and PBIC of the restricted models with equal threshold means across countries ($\mu_1 = \mu_2 = \dots = \mu_C = 0$) are 335,584 and 334,972, respectively. For the model with no within-country threshold heterogeneity ($v_1 = v_2 = \dots = v_C = 0$), these statistics are equal to 339,379 and 338,768, respectively. These information criteria are larger than those of the unrestricted model (335,494 and 334,872). This indicates that there are differences in response tendencies between, but also within countries and supports the need for accommodating heterogeneity in excess of that described by latent segments in standard mixture models.

Figures 4.1 to 4.4 show the MEC structures for the segments derived from the link strengths, represented graphically as 'means-end maps.' The attributes are located at the bottom of the figures (in arbitrary order), the benefits in the middle and the values at the top. (Note that the

concepts are abbreviated; see Appendix 4.A for the complete formulations.) For each segment, lines represent the major links between attributes and benefits (AB-links), and benefits and values (BV-links). In addition, for each link we present $p_{js}^m = \exp(u_{js}^m) / [1 + \exp(u_{js}^m)]$, the segment-specific probability that a link is chosen, adjusted for response tendencies (p_{js}^m is obtained from Equation (4.1), with $\theta_i = 0$ for all i and u_{js}^m substituted for u_j^m). To identify the major links, we tested if each link's strength u_{js}^m is significantly larger than a chosen value θ^* . The lower the value of θ^* , the more of the u_{js}^m will exceed the subject-specific thresholds. To provide a parsimonious representation and avoid clutter of the graphs, we chose $\theta^* = -.8$ as an appropriate cut-off level, corresponding to a minimal link probability of about .3. In the figures, the thickness of the lines reflects the corresponding link probability.

On examining the four MEC maps simultaneously, a number of APT links appear to be important across all four segments. Those links are important to all consumers in the 11 countries. Means-end chains that appear consistently across segments present opportunities for the development of standardized products, supported by mass (unsegmented) communication. Of particular interest are complete chains, or unique perceptual orientations, of a particular attribute, benefit, and value, which present product positioning options. The chain 'low-fat – good for health – fun and enjoyment in life' is found in all four segments (although the connections are considerably stronger in *S3* than in *S1* and *S4*). This suggests opportunities for low-fat yogurt products that are mass-marketed to the EU-consumers in general, with the resulting fun and enjoyment from good health as the consumer's driving value orientation in the communication campaign. The BV-link between 'good for your health' and 'fun and enjoyment' is common across segments. However, major differences between the segments are also present.

Figure 4.1
Probabilistic means-end map segment S1

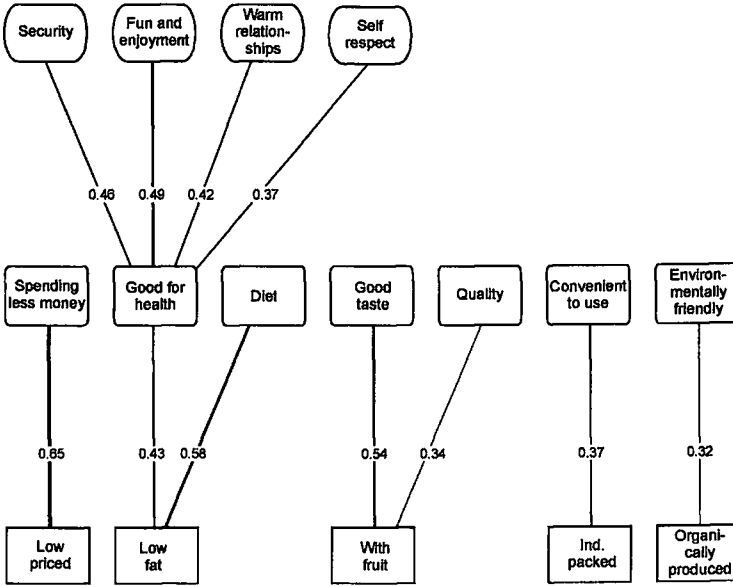


Figure 4.2
Probabilistic means-end map segment S2

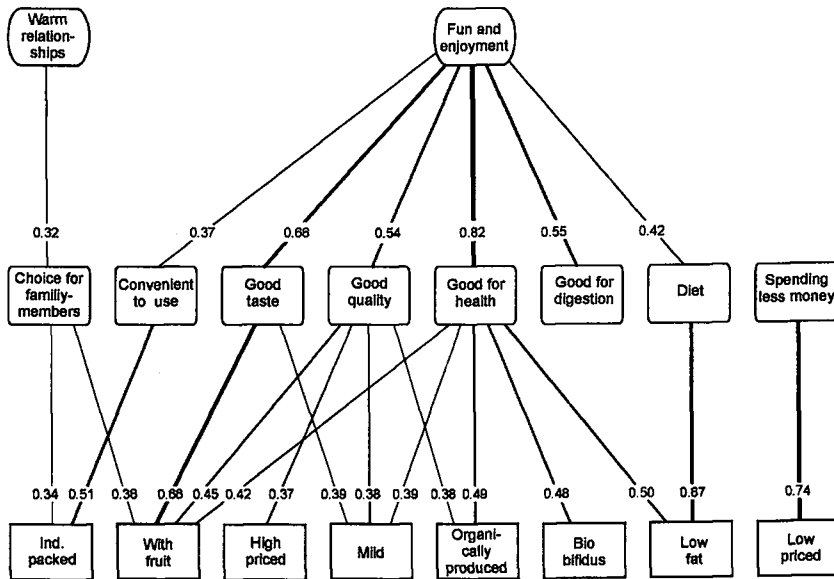


Figure 4.1 displays the main links in the means-end structure for segment *S1*. Consumers in this segment do not have as rich cognitive structures concerning yogurt as those in the other segments. Seven AB-links and four BV-links appear in the map. Each attribute leads to one or two perceived benefits, but only one benefit connects to the value level. The central benefit is 'good for your health,' which is linked with four values. The attribute 'low fat' leads to this central benefit. The four values present different positioning options by associating healthy, low-fat yogurt with security, hedonic, social, or self-respect values in advertising campaigns.

The MEC-map of segment *S2* in Figure 4.2 displays a rich structure of links: 16 AB-links and 7 BV-links arise. Multiple links connect to 'good for your health' and 'good quality,' which seem to be central benefits for this segment. Subjects in this segment seem to categorize yogurt at a higher level of abstraction than subjects in segment *S1*, viz. the value-level. 'Fun and enjoyment in life' is clearly the need that this segment seeks to fulfill through the consumption of yogurt. Many benefits contribute to fulfillment of the need for fun and enjoyment, but especially 'good taste' and 'good for your health'. A dominant MEC chain is the link between 'with fruit,' 'good taste,' and 'fun and enjoyment in life,' both link probabilities being .66. In developing and communicating yogurt for this segment, fun and enjoyment in life should provide the major perceptual orientation. The multiple links with benefits, in particular with taste, quality, and health and the large number of product attributes providing those benefits provide many opportunities for product development and positioning.

Consumers in segment *S3* have many strong links, especially at the AB level. The MEC-map in Figure 4.3 consists of 29 AB and 10 BV links. Consumers' cognitive structure with respect to yogurt appears to be more elaborate than in the two previous segments, but at the same time seems to lack a single benefit or value that is central in the cognitive associative network. This suggests that multiple product options may be of interest to this segment. At the attribute level, 'organically produced,' 'bio-bifidus,' 'low fat,' and 'with fruit' have many strongly associated

benefits. In this segment, consumers derive health benefits from organically produced, low-fat yogurt, and yogurt with bio-bifidus. The health benefits 'good for health,' 'good for digestion,' and 'replaces unhealthy snacks' have particularly rich associations with the attributes. Segment *S3* shows the richest complete means-end chains involving diet. Low-fat contributes to the product being 'good if one is on a diet,' which in its turn contributes to the three internally-focused values (Homer and Kahle 1988) 'self-respect,' 'sense of accomplishment,' and 'self-fulfillment'. To compare, in *S1* and *S4*, diet is not related to any value while in *S2*, it is related to the interpersonal value (Homer and Kahle 1988) 'fun and enjoyment'.

Segment *S4* contains fewer links than segments *S2* and *S3*, but its cognitive structure is somewhat richer than that of segment *S1*. In the map 12 AB-links and 5 BV-links are present (see Figure 4.4). As compared to the other segments, one of the distinguishing characteristics of this segment is the link between 'good quality' and 'security.' Whereas the link probability is smaller than some of the other link probabilities, this link does not arise in any of the other segments. The attributes 'high priced' and 'organically produced' are linked to 'good quality,' forming a distinguishing MEC network that can serve as the basis for development of higher priced organic quality yogurts. In advertising campaigns, this product can be linked to security needs.

Summarizing, the results show distinct and actionable segments. In the cross-national segment *S1*, consumers are typically focused on the links between health benefits and various values, whereas the consumers in segment *S2* use multiple benefits for satisfaction of fun and enjoyment in life. The North-European segment, *S3*, typically derives a broad set of health-related benefits from multiple yogurt attributes, offering various (product development) means to the same benefit. The pan-European segment *S4* displays much of the common MEC structure across segments – which is not surprising as it is the largest segment – but it has a distinguishing MEC component that centers around product quality.

In the remainder of this section, we will construct profiles of the segments by relating them to descriptive consumer data. Then, we assess

the predictive validity and demonstrate the incremental contributions of our model over existing methods.

Figure 4.3
Probabilistic means-end map segment S3

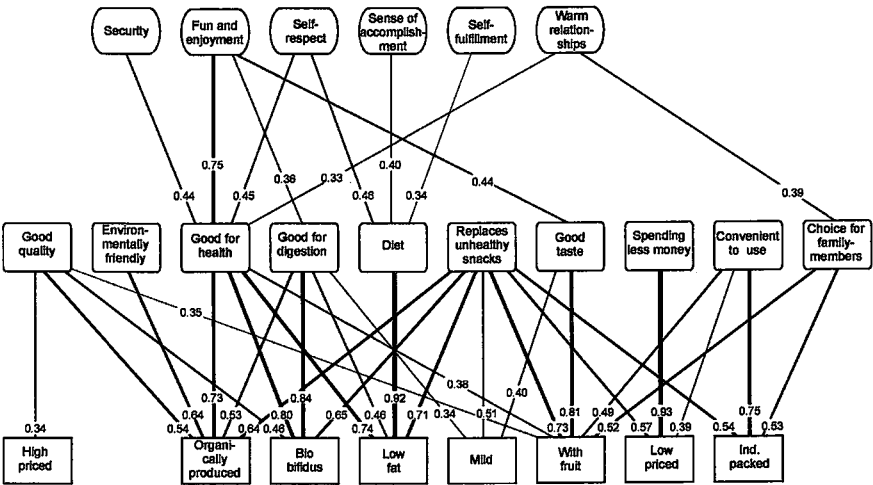
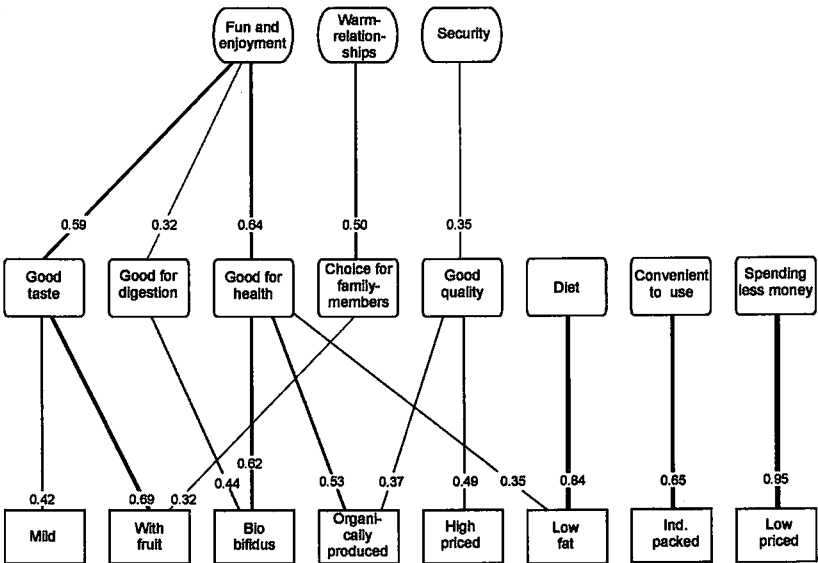


Figure 4.4
Probabilistic means-end map segment S4



4.5.3 *Segment Profiles*

The segments were related to descriptor data, including sociodemographics, consumption patterns, information on media consumption, and information on personality and attitudes. The personality and attitude constructs were: deal proneness (cf., Lichtenstein, Netemeyer, and Burton 1995), consumer innovativeness (Steenkamp, ter Hofstede, and Wedel 1999), consumer ethnocentrism (Shimp and Sharma 1987), and product involvement (Mittal and Lee 1989). We constructed profiles of the four segments using the regression-based procedure of Kamakura and Mazzon (1991). The results, based on a 5% level of significance, are presented in Table 4.6.

Consumers in Segment *S1* tend to be older, less educated, have lower incomes, and live in less urbanized areas. They spend less on yogurt and tend to visit smaller stores more often for their yogurt purchases. They tend to listen more often to the radio, watch more serials and entertainment programs, are more ethnocentric, and tend to be less innovative, deal prone, and involved with yogurt. The socio-demographic profile of *S2* is similar to *S1*. Further, these consumers use yogurt more often as a snack and tend to purchase it more often in larger stores. They tend to be more inclined to respond to promotions, less innovative and more ethnocentric. Their attitudes and involvement regarding yogurt are generally higher. Segment *S3* consists of younger consumers that are higher educated and have higher incomes. They are relatively light users of yogurt, but use it more often as a snack. These consumers are more exposed to radio and newspaper advertising. They are more deal prone and innovative, less ethnocentric and have a more positive attitude toward yogurt in general.

Consumers in Segment *S4* have a socio-demographic profile similar to segment *S3*, but live more often in urban areas. They spend most on yogurt, which they purchase more often in larger stores. Except for movies, these consumers are generally less exposed to media. In sum, the significant effects show a coherent pattern and contribute to the identifiability and accessibility of the segments.

Table 4.6
Segment profiles

<i>Dependent variable</i>	<i>R^a</i>	<i>Segment^b</i>				
		<i>S1</i>	<i>S2</i>	<i>S3</i>	<i>S4</i>	
<u><i>Socio-demographics</i></u>						
Age	.114	+	+	−	−	
Highest level of education	.148	−	−	+	+	
Total household income after tax (in ECU)	.138	−	−	+	+	
Size of the place of residence	.110	−			+	
<u><i>Consumption patterns</i></u>						
Expenditures on yogurt (in ECU)	.089	−		−	+	
Frequency of using yogurt as a snack	.110	−	+	+		
Purchase frequency in convenience stores	.077	+			−	
Purchase frequency in specialty stores	.077	+	−			
Purchase frequency in hypermarkets	.118	−	+	−	+	
<u><i>Media consumption</i></u>						
Frequency of listening to the radio	.105	+		+	−	
Frequency of reading daily newspaper	.114		+	+	−	
Frequency of watching serials	.152	+			−	
Frequency of watching entertainment programs	.184	+			−	
Frequency of watching movies	.084			−	+	
<u><i>Personality and attitudes</i></u>						
Deal proneness	.122	−	+	+		
Consumer innovativeness	.145		−	+		
Consumer ethnocentrism	.145	+	+	−		
Involvement with yogurt	.148	−	+		+	
Overall attitude toward yogurt	.122	−	+	+		

a All R-values are significant at the .01-level.

b The plus and minus signs reflect the direction of the significant coefficients (at the 5% level) of regressing the dependent variables on the posterior segment memberships. The intercept was restricted to the mean value of the dependent variable. Hence, a positive (negative) sign indicates that the behavior or attitude of consumers within a segment is above (below) average.

4.5.4 Predictive Validity of the Model

In order to assess the predictive validity of our model, we cross-validated the results using holdout predictions of segment memberships. After assigning all respondents to the segment with largest posterior segment membership probability, estimation and holdout samples were constructed and discriminant analysis was applied to the estimation samples. The discriminant function was constructed from the APT links to explain segment-membership. The segment memberships in the holdout samples were then predicted using the discriminant function.

In predicting the segment memberships, two approaches have been taken. First, estimation and holdout samples were constructed by assigning at random two-thirds of the subjects to the estimation sample and one-third to the holdout sample. Second, we used the U-method (cf., Ghose 1998), which is closely related to the jackknifing method of classification. The U-method creates holdout samples of size one, so that holdout predictions are obtained for each individual separately. This procedure repeatedly omits one observation in estimation of the discriminant model and then predicts the segment membership of that omitted observation using the estimated discriminant function.

The procedures produced similar results. The first approach resulted in a prediction accuracy of 86 percent, ranging from 82 percent in segment *S2* to 88 percent in *S4*. Thus, 86 percent of the cases is classified correctly, and no segment suffers from misclassification. The U-method resulted in a prediction accuracy of 87 percent, ranging from 82 percent in segment *S2* to 91 percent in *S4*. The high prediction accuracy strongly supports the stability and predictive validity of the MEC segment structure identified with our model.

4.5.5 Comparison with Other Methods

Two alternatives for analyzing the pick-any data obtained from the APT-tasks are correspondence analysis and cluster analysis. Correspondence analysis relates to the nature of the binary APT data, but does not arrive at the objective of our study, i.e., identifying segments of consumers with similar means-end chains and construction of the MEC maps. Rather it

identifies continuous dimensions of subjects and concepts underlying the APT matrices. Cluster analysis, and in particular K-means, is the prevailing approach to identify international segments (e.g., Boote 1983; Kale 1995; Yavas, Verhage, and Green 1992). One of the best clustering methods is K-means (Wedel and Kamakura 1998). To assess if and to what extent our model outperforms such a segmentation approach, we compare our model to K-means clustering.

Our model has several important features that are not accommodated by cluster analysis. First, a probabilistic classification method is used, as opposed to deterministic classification with K-means. Second, the model captures heterogeneity in response behavior, both within and between countries. Third, our model explicitly captures differences in segment sizes between countries, whereas K-means does not. Finally, the pseudo likelihood is maximized, which accommodates the international sampling design and therefore leads to unbiased parameter estimates. K-means maximizes the ratio of the within-cluster and the between-cluster variance, and does not account for the sampling design. It may be expected that these features increase the performance of our model as compared to K-means.

We used the pan-European data set to compare our model to a K-means clustering solution. Three criteria for model comparison are used: model fit in terms of pseudo likelihood, correspondence of segment memberships, and the actionability of the segments. First, the fit of our model exceeds the fit of the K-means clustering solution for four segments. The difference in log-pseudo likelihood was a substantial 4691.12.

Second, the correspondence between the two solutions was rather low. We cross-classified the segments identified by the two methods.³ The percentage of respondents that were classified by both approaches to the same segment (after appropriate permutation) ranged from about 40 percent for segment 1 to about 70 percent for segment 4 and was on average below 50 percent. This is confirmed by the Rand index (Rand

³ For our model, the subjects were classified according to the maximum a-posteriori segment memberships.

1971), $RI=.597$, which indicates that the probability that the methods allocate pairs of subjects in a similar way is around 60 percent. Note that this level of agreement is only slightly higher than that obtained when our method is compared with a random assignment of subjects ($RI=.584$).

Third, K-means clustering identified segments that were much less actionable. These segments discriminated mainly on the overall number of links that were chosen. We counted for each segment the number of times a link has the highest probability compared to the other segments. For K-means we arrive at 4, 82, 15, and 0 percent of the links having maximal probability in clusters one to four, respectively. For the solution of our model, these figures were 18, 32, 41, and 9 percent. The more even distribution of links indicates that our model resulted in more actionable segments. Each segment has its own idiosyncratic pattern of links, while in the K-means-solution one segment (segment S2) in particular stands out because it provided an excessive number of dominant links. Apparently, these disappointing results of K-means are due to the fact that K-means does not account for differences in response behavior and hence leads to substantially different results.

4.6 Discussion

This chapter proposed an approach to international market segmentation using means-end chains, which integrates international market segmentation and product positioning. Using MEC as segmentation basis combines the advantage of product-specific and consumer-specific bases. It supports product positioning not only at the level of product attributes, but also includes the benefits derived by segments from those attributes, as well as the values satisfied. This increases the actionability of the identified segments for targeted strategies of international product development and communication. The proposed model provides segment-specific estimates of strengths of links that tie the consumer to the product, i.e., that identify cognitive associations between product attributes, benefits of product use, and consumer values at the segment-level. It explicitly captures within- and between-country heterogeneity in

response behavior. Heterogeneity was found to be highly significant both within and between countries. The model is tailored to binary MEC data collected by APT, validated in chapter 3, and successfully applies the PML approach proposed in chapter 2 to account for the international sampling design.

We applied the methodology to identify segments in the European yogurt market, using a large cross-national sample of European consumers. Four segments were identified, of which one was truly pan-European. The segments were found to be actionable and the pattern of links between attributes and values gave rise to strategic implications with respect to product development and communication. In addition, consumer characteristics were found to be related to segment membership, which contributes to the identifiability and accessibility of the segments. We also found high predictive accuracy in holdout samples and an empirical comparison with K-means clustering demonstrated the better performance of our model.

While our model was developed for segmentation in international markets, it can easily be adapted to the case of domestic segmentation, i.e., segmentation within a single country. Our model nests a domestic segmentation model, which is obtained by imposing the following restrictions on the model: $\mu_c=0$, $v_c=v$ and $\pi_{s|c}=\pi_s$ for all $c=1,...,C$. These restrictions come down to dropping the mean threshold parameters in Equations (4.3) and (4.4) and concomitant variables in Equation (4.5), and estimating a single variance parameter for the thresholds in Equations (4.3) and (4.4). The log-pseudo likelihood in Equation (4.6) of this chapter reduces to the standard log-likelihood: $\ln L_p = \sum_i \ln L_i$, and the PML approach reduces to a standard maximum likelihood approach.

For international marketing purposes, the methodology lends itself quite well to different types of international target market selection and differentiation strategies (which are not mutually exclusive; cf., Kotler 1997). A first option is to develop specific products for specific segments. Second, a single segment may be targeted by a bundle of products (market specialization) to preempt competitive efforts and to capture the variety seeking tendencies of consumers (Baumgartner and Steenkamp

1996). Third, basically the same physical product can be developed for multiple segments. However, the product should be positioned differently in these segments based on the segment-specific benefits and values to which the product's main attributes are perceived to lead. Such a product specialization strategy will result in substantial cost reductions in production. A fourth strategy is to develop a mass product, based on communalities in MECs between segments (if any). Further, our approach provides guidelines for integrated communication strategies by revealing meaningful linkages between means-end chain elements (Reynolds, Gengler, and Howard 1995).

Our MEC-based approach is not only useful for product positioning, but also as input for the new product development process (NPD), which involves a range of decisions from the identification of market opportunities to market forecast and communication. Our MEC-based approach is obviously not the only input to this complex managerial process involving marketers, R&D personnel, production managers and so on, but it may be an important one. A key aspect of NPD is the core benefit proposition: a short list of the strategic benefits that the new product provides to its customer segment, and clues of how the product provides these benefits (Urban and Hauser 1993). Our procedure produces that information at a segment level. Obviously this information should be supplemented by for example, cost and competitive information. Our methodology, being based on MEC, is also a useful tool for coordinating thinking and communication across functional areas, identified by Griffin (1992) as a major way to improve the NPD process.

Our analysis has been restricted to a single product category. Further research may focus on other categories and/or branded products to assess the pan-regional or global nature of the segments that our approach identifies. The inclusion of brands can be accommodated through an additional brand-attribute matrix in the sequential APT task (cf., Gutman 1982). However, it needs first to be assessed whether the brand perceptions in the brands-attributes matrix are independent of the links between attributes, benefits, and values. Loglinear models similar to those formulated in chapter 3 may be specified that allow to test such

independence relations. If the relations turn out to be independent, the brand-attribute links may be added to the model. This is an issue for further research.

Another limitation of our empirical application is that our data were restricted to a limited set of countries, i.e., 11 countries within the EU. Whereas there is quite some cultural diversity between the European countries, the results cannot be generalized directly to other countries and do not provide an answer to the existence of global segments covering different parts of the world (e.g., Europe, Asia, Africa). Further research may focus on this issue and apply the methodology proposed in our study to countries from different continents.

We believe that the proposed methodology, integrating international segmentation theory, measurement, and analysis in one framework, provides a step toward a closer fit of marketing research and managerial objectives in the international marketing domain.

Appendix 4.A

In APT, respondents are confronted with two matrices, an AB and BV matrix. The matrices that have been used for the purpose of this study are given in Figure 4.A. In the AB matrix, a priori defined attributes and benefits comprise the rows and columns, respectively. Similarly, the BV matrix includes benefits and values as rows and columns. For each column in the AB matrix (BV matrix), respondents indicate the benefits (values), if any, perceived to be associated with each of the product attributes (benefits). The validity of APT to collect data on consumer means-end chains has been demonstrated in chapter 3.

Attribute-Benefit Matrix

[illegible]

Benefit-Value Matrix

with:

I associate ...

[illegible]

Appendix 4.B

Following Bock and Lieberman (1970), the density unconditional on θ in Equation (4.3) may be approximated numerically by:

$$(4.B_1) \quad f_{i|s}(y_i; u_s, \mu_c, v_c) = \frac{1}{\sqrt{2\pi}} \sum_{r=1}^R a_r \exp\left(\frac{1}{2} g_r^2\right) f_{i|s}(y_i | g_r; u_s, \mu_c, v_c),$$

where a_r and g_r denote the abscissae and respective weights of the tables for the R -th order Gauss-Hermite quadrature. Substitution of (4.B₁) in the log-pseudo likelihood in Equation (4.6) yields:

$$(4.B_2) \quad \ln L_P = \sum_{i=1}^I w_i \sum_{s=1}^S \pi_{s|c(i)} \sum_{r=1}^R \tilde{a}_r f_{i|s}(y_i | g_r; u_s, \mu_{c(i)}, v_{c(i)}),$$

where $\tilde{a}_r = a_r \exp(\frac{1}{2} g_r^2)$. The derivative of the log-pseudo likelihood with respect to γ_{sc} is equal to

$$(4.B_3) \quad \frac{\partial \ln L_P}{\partial \gamma_{sc}} = \sum_{i:c(i)=c} w_i (\alpha_{is} - \pi_{s|c}),$$

where the α_{is} are the posterior probabilities obtained by Bayes rule. In order to assure the threshold variances to be positive, we use a reparameterization: $v_c = \exp(\tau_c)$. The derivatives of the log-pseudo likelihood with respect to u_{js}^m , μ_c , and τ_c are:

$$(4.B_4) \quad \frac{\partial L_P}{\partial u_{js}^m} = \sum_{i=1}^I w_i \frac{\pi_{s|c(i)}}{f_i(y_i; u, \mu_{c(i)}, v_{c(i)})} \sum_{r=1}^R \tilde{a}_r f_{ir|s}(y_i | u_s, \mu_{c(i)}, v_{c(i)}) \\ \times \left(y_{ij}^m - \frac{\exp(u_{js}^m - \exp(\tau_{c(i)}) g_r - \mu_{c(i)})}{1 + \exp(u_{js}^m - \exp(\tau_{c(i)}) g_r - \mu_{c(i)})} \right)$$

$$(4.B_5) \quad \frac{\partial L_P}{\partial \mu_c} = - \sum_{i:c(i)=c} \frac{w_i}{f_i(y_i; u, \mu_c, v_c)} \sum_{s=1}^S \pi_{s|c} \sum_{r=1}^R \tilde{a}_r f_{ir|s}(y_i | u_s, \mu_c, v_c) \\ \times \sum_{m=1}^2 \sum_{j=1}^{J_m} \left(y_{ij}^m - \frac{\exp(u_{js}^m - \exp(\tau_c) g_r - \mu_c)}{1 + \exp(u_{js}^m - \exp(\tau_c) g_r - \mu_c)} \right)$$

$$\begin{aligned}
 (4.B_6) \quad \frac{\partial L_p}{\partial \tau_c} = & - \sum_{i:c(i)=c} w_i \frac{\exp(\tau_c)}{f_i(y_i; u, \mu_c, v_c)} \sum_{s=1}^S \pi_{s|c} \sum_{r=1}^R \tilde{a}_r f_{ir|s}(y_i | u_s, \mu_c, v_c) \\
 & \times \sum_{m=1}^2 \sum_{j=1}^{J_m} \left(y_{ij}^m - \frac{\exp(u_{js}^m - \exp(\tau_c) g_r - \mu_c)}{1 + \exp(u_{js}^m - \exp(\tau_c) g_r - \mu_c)} \right)
 \end{aligned}$$

The pseudo maximum likelihood estimates are obtained by equating expressions in Equations (4.B₄), (4.B₅), and (4.B₆) to zero and solving for the unknown parameters.

Chapter 5

A Bayesian Approach to Identifying Geographic Target Markets

5.1 Introduction

In response to the increasing globalization of the business world, companies are competing for a presence in different markets around the world. Faced with an increasing saturation in home markets and lured by growth opportunities, many companies are forced to extend their business internationally. Once a company commits itself to entering into foreign markets, management is confronted with the task of emphasizing activity geographically and of determining the right geographic markets to target. The key to success is to understand the preferences and behavior of consumers in these target markets, and tailor strategies to their needs (Douglas and Craig 1992; Jain 1990). So, a central issue for international marketers is how to determine the geographic location of such target markets and to design marketing strategies that tap the needs of consumers in these areas.

Several approaches are conceivable to identify geographic target segments. A natural way to define such segments is to use the existing national or political borders. Such markets are typically reflected by countries, trade areas such as the European Union, or regions within countries. A rationale for entering such areas is that it results in *accessible* and *costs effective* strategies through centralization of activities such as production, sales force management, service support, logistics, advertising, production of promotional materials, and training of personnel (Porter 1998; Takeuchi and Porter 1986; Yip 1995). Physical distribution of products and services is more efficient if a company's activities are located nearby (Douglas and Wind 1987). Customer support for handling guarantees and customer service places heavy demands on companies, depending on the geographic location of markets (Takeuchi and Porter 1986). In addition, advertising and promotional activities are frequently differentiated across regions. For example, Campbell soup has created different product versions and sales promotions for different regions in the United States (Wilkie 1990), and much stronger regional diversification strategies are pursued by food companies such as Unilever. Clearly, the advantages of such strategies concentrate around accessibility and efficiency.

As opposed to defining target segments based on political borders, several studies propose to group countries or regions using their characteristics (Askegaard and Madsen 1998; Helsen, Jedidi and DeSarbo 1993; Kahle 1995; Sethi 1971). Different variables have been proposed as international segmentation basis, such as socioeconomics, cultural variables, and consumption patterns. This allows firms to target countries or regions belonging to the same segment with similar marketing strategies. If response-based segmentation is used the segments will in addition be more responsive to marketing effort. After all, the segments are constructed on the basis of the similarity of variables relevant for formulating marketing strategies. However, whereas such segments may be more *responsive*, their typical geographic configuration is often less *accessible* and in some cases even renders marketing and logistic operations infeasible. For example, entering a segment consisting of

consumers in Spain, Uruguay, and Israel with a single perishable product will often lead to operations that are too costly.

It is critical that a geographic segment defines a single contiguous area to warrant accessibility and cost effectiveness of marketing strategies. For example, in distribution intensive industries such as retailing and in industries dealing with perishable products, geographically dispersed target segments will often not allow profitable entry strategies to be pursued. One example of such region-based entry strategy is that of “staggered regional” rollout, where a company enters a new target market in a convenient (sub)region and gradually extends its marketing and sales activities across boundaries into contingent areas.

In summary, there is a clear trade-off between *accessibility* and *responsiveness*. Whereas the accessibility of geographic location poses restrictions on the configuration of target areas, identification of target segments should be based on consumer needs for formulating responsive marketing strategies. This calls for a methodology that identifies spatial segments, based on the similarity of consumers’ needs and that gainfully exploits geographic contiguity constraints. In this chapter, we seek to extend previous approaches by developing a model that identifies responsive geographic segments based on consumer needs, but at the same time enforces geographic contiguity to yield accessible and cost-effective segments.

We develop a flexible model that identifies geographic target segments that obey contiguity constraints. In the interest of the responsiveness of the geographic segments, the model uses a consumer-level segmentation basis and captures heterogeneity that is likely to exist within geographic segments. In specifying the model, we take a hierarchical Bayes approach that accommodates a broad set of restrictions on segments. To assess the performance of the model we conduct a synthetic data analysis of eight data sets that are generated according to different experimental conditions. We apply the methodology in the international retailing domain, where international expansion is currently a major growth strategy (Bell, Davies, and Howard 1997). On the basis of the importance that consumers attach to different attributes of store image, we identify geographic target markets in the European Union. The

segments are related to relevant characteristics and the results are empirically compared with an unrestricted solution of the model. Finally, we discuss our findings and conclude with issues for further research.

5.2 Model Development and Estimation

5.2.1 Previous Approaches to Constrained Segmentation

Our objective is to develop a segmentation model that divides an area (e.g., the United States or the European Union) in a number of unobserved contiguous geographic segments, each representing a potential target market to enter. These segments are groupings of predefined regions (e.g., states or counties in the United States, Länder in Germany, départements in France). The regions are grouped into geographic segments based on observations of consumers in these regions. Basically, this problem reduces to a constrained segmentation of regions, where particular restrictions are imposed on the segment memberships of these regions.

Several studies have looked into the issue of constrained segmentation. Broadly two different classes of restrictions are distinguished. First, constraints may be imposed on the segmentation basis itself. For example, segment-specific parameters may be subjected to non-negativity constraints to ensure positive price elasticities within segments. Such constraints have been denoted as *internal constraints*. On the other hand, *external constraints* impose restrictions on the segmentation entities, i.e., the regions that are aggregated into segments. External constraints define a restriction on the number of partitions of the set of segmentation entities and are independent of the segmentation basis. The contiguity and other restrictions considered in this study belong to the latter type of constraints.

In the classification literature, several studies have considered the problem of segmentation with external constraints. Gordon (1973) developed a clustering algorithm that identifies clusters obeying contiguity constraints in a single dimensional space where the units in the

segmentation solution are ordered (e.g., the time dimension). The method identifies the location of a barrier between neighboring objects based on the similarity of the set of objects on either side of it. Ferligoj and Batagelj (1982) present hierarchical clustering and a local optimization procedures that accommodate all sorts of symmetric relational constraints, i.e., constraints that are defined on symmetric relations between the segmentation objects. The authors apply the method to identify contiguous segments of countries based on developmental indicators, where the relations between the countries are defined by their geographic adjacency. For hierarchical clustering, they extend the methodology to accommodate asymmetric relations (Ferligoj and Batagelj 1983). DeSarbo and Mahajan's (1984) CONCLUS algorithm extends these clustering techniques to accommodate an even broader class of both internal and external restrictions. The approach allows overlapping clusters and restrictions on the compactness of the segments. The algorithm is applied to a sales territory delineation problem to ensure equal workload and sales potential for each salesman.

Whereas the above mentioned studies enable the identification of contiguous geographic segments, they all rely on cluster analysis. Cluster analysis has several disadvantages. It is based on mostly arbitrary distance measures and does not provide reliability measures of the results. In addition, these procedures are limited to input data that are given at the same level as the units to be classified, i.e., the regions. If multiple observations are obtained per region, sequential procedures are needed, where region-specific estimates are identified in a first stage and cluster analyzed in a second stage. Such approaches are cumbersome because they do not take the unreliability of the region-specific estimates into account. Moreover, estimates may not be identified for regions with limited observations (e.g., regression parameters may not be identified due to insufficient degrees of freedom in a region).

In response to the above-mentioned limitations of cluster analysis, mixture models have been proposed to identify segments (e.g., Green, Carmone, and Wachspress 1976; Wedel et al. 1995). Mixture models treat the segmentation problem in a probabilistic context and provide

standard errors of parameters. They provide a model-based approach to segmentation. However, whereas internal constraints are feasible (e.g., DeSarbo and Edwards 1996), imposing restrictions on the segmentation units (such as contiguity) involves complex integrals for which no closed form expressions can be derived. Moreover, accounting for heterogeneity within segments is not easily accommodated in the standard ML-based mixture framework. In chapter 4, for example, we accounted for heterogeneity of a single parameter, which involved complex integration procedures. In higher dimensions the complexity of such procedures will increase substantially.

Instead, we will formulate a Bayesian segmentation methodology that accommodates external restrictions on the segments and treats the problem in a probabilistic context. We develop the methodology for the case of response-based segmentation, which has seen broad application in market segmentation. The approach overcomes the above mentioned problems of cluster analysis and ML-based mixture models. We build upon the Bayesian hierarchical mixture model (e.g., Allenby, Arora, and Ginter 1998) to accommodate external restrictions on the segments. In this formulation, the model will take both the differences within and between segments into account.

5.2.2 A Response-based Geographic Segmentation Model

In this section, we develop a geographic segmentation methodology for response-based segmentation. We consider the case of a multi-attribute model, where evaluations of objects are related to attributes of those objects. Such an approach covers numerous applications in market segmentation and allows the researcher to accommodate most (unobserved) product-specific segmentation bases, including benefits, product perceptions, and product attribute importance. The formulation enables the modeling of many applications in marketing including metric conjoint analysis (Vriens, Wedel, and Wilms 1996), the formation of store image (Steenkamp and Wedel 1991), trade show performance evaluation (DeSarbo and Cron 1988), and customer satisfaction (Wedel and DeSarbo 1995). In the discussion section of this chapter, a

generalization of the model is provided that accommodates other segmentation bases.

In the model we distinguish three levels. At the highest level, geographic segments need to be identified that are not observed a priori. At the second level, these segments consist of predefined regions, the basic building blocks that constitute the geographic segments. The lowest level of the model is defined by consumers within regions. Let:

$s=1,...,S$ unobserved geographic segments,

$r=1,...,R$ predefined regions,

$i=1,...,I_r$ subjects in region r .

$k=1,...,K$ attributes

$l=1,...,L$ objects

In this formulation, subjects are nested within regions, and predefined regions are nested within segments. The multi-attribute model assumes that the overall evaluation y_{ril} of a particular object l (e.g., a store) can be expressed as linear combinations of perceptions of that object's attributes (e.g., a store's image attributes) x_{rik} . Therefore, at the lowest level of the model, we obtain a standard linear regression equation, as follows:

$$(5.1) \quad y_{ril} = \sum_{k=1}^K x_{rik} \beta_{rk} + \varepsilon_{ril},$$

where β_{rk} denotes the importance of attribute k in region r . Note that we assume the attribute importances to vary across regions. The disturbances have independent normal distributions with zero mean and variance σ_y^2 :

$$(5.2) \quad \varepsilon_{ril} \sim N(0, \sigma_y^2).$$

A sequential approach would estimate the β_{rk} for each region r ($r=1,...,R$) separately, using Equations (5.1) and (5.2) and cluster analyze the parameter estimates in a second stage. However, to accommodate differences of the β_{rk} across consumers within regions, we build a hierarchical structure on top of Equation (5.1). At the regional level of the model, we add a data augmentation step (Tanner and Wong 1987) by introducing unobserved segment memberships indicators of regions, ξ_r ,

(cf., Robert 1996). We let the β_{rk} follow normal distributions within segments, as follows:

$$(5.3) \quad [\beta_{rk} | \xi_r = s, \bar{\beta}_{sk}, \sigma_k^2] \sim N(\bar{\beta}_{sk}, \sigma_k^2),$$

with ξ_r the unobserved segment membership of region r , $\xi_r \in \{1, \dots, S\}$. In this formulation β_{rk} has a segment-specific mean $\bar{\beta}_{sk}$ and variance σ_k^2 , conditional on region r belonging to geographic segment s .

The $\bar{\beta}_{sk}$ reflect the segment-specific importance of the attributes and indicate which image attributes need to be emphasized, given that certain positions are not taken by competitors. The σ_k^2 account for the within-segment heterogeneity, i.e., the differences in β_{rk} that exist within segments across consumers. Note that both $\bar{\beta}_{sk}$ and σ_k^2 are indicators of responsiveness of segments. The differential responsiveness will improve with increasing differences of the $\bar{\beta}_{sk}$ among segments, and with diminishing values of σ_k^2 . So, ideally, we would identify segments that have large variation in the pattern of the $\bar{\beta}_{sk}$ across segments, and small values of σ_k^2 .

Equation (5.2) is defined conditionally on a particular partition of the regions into segments. Let $\Xi = (\xi_1, \dots, \xi_R)$ a vector of segment memberships. Ξ defines a partition of the R regions into S segments. The contiguity restrictions depend on the geographic pattern of regions, contained in an $(R \times R)$ -matrix $G = [g_{rr'}]$. Let $G_s(\Xi)$ be a matrix consisting of the columns and rows of G , that correspond to the regions belonging to segment s . Segment s is contiguous if the Markov chain defined by $G_s(\Xi)$ is *connected*, i.e., each state of the chain can be reached from another state in a finite number of transitions.

In Markov chain theory, the matrix $(G_s(\Xi))^j$ for a fixed value of j has a special interpretation. Its entries provide information about the *j-step transitions* between states (Winston 1987). An entry $(G_s(\Xi))_{rr'}^j$ indicates the number of paths from region r to region r' , that run through

the regions in segment s in j steps. If a segment is contiguous, then in the worst case, a path runs through all other regions in that segment. Hence, two regions are connected if and only if there is a nonzero j -step transition between those regions, with j less than the number of regions in segment s , which means that segment s is contiguous if $\sum_{r=1}^{N_s-1} (G_s(\Xi))^r > 0$, with $N_s = \sum_{r=1}^R I(\xi_r = s)$ the number of regions in segment s . Given G , the set of partitions that satisfy the contiguity restrictions is defined as:

$$(5.4) \quad \Omega = \left\{ \Xi \left| \sum_{j=1}^{N_s-1} (G_s(\Xi))^j > 0, \forall s = 1, \dots, S \right. \right\}.$$

We need to define the probability distribution conditional on the feasible set of partitions contained in Ω . However, a general expression for the probability of a partition will generally not exist in closed form. Therefore, we define a Gibbs field on the regions (cf., Geman and Geman 1984), which allows treating each region's prior segment membership probability separately by conditioning on the segment memberships of the other regions, as follows:

$$(5.5) \quad p(\xi_r = s | \Xi_{/r}, \Xi \in \Omega) = \frac{\pi_{s0}}{\sum_{s' \text{ adjacent to } s} \pi_{s'0}},$$

with π_{s0} the (prior) segment size of segment s and $\Xi_{/r}$ contains the indicators in Ξ excluding ξ_r . In this formulation, the probability of an individual region belonging to a particular segment is a function of the segment sizes bordering that region. Thus, a multinomial distribution applies to ξ_r given $\Xi_{/r}$ with probability as given in Equation (5.5).

To finalize the hierarchical Bayes model, we specify hyper distributions for σ_y^2 , σ_k^2 , and $\bar{\beta}_{sk}$ as follows:

$$(5.6) \quad \sigma_y^{-2} \sim \text{Gamma}(\alpha_{y0}, \delta_{y0}),$$

$$(5.7) \quad \bar{\beta}_{sk} \sim N(\mu_{\beta0}, \sigma_{\beta0}^2),$$

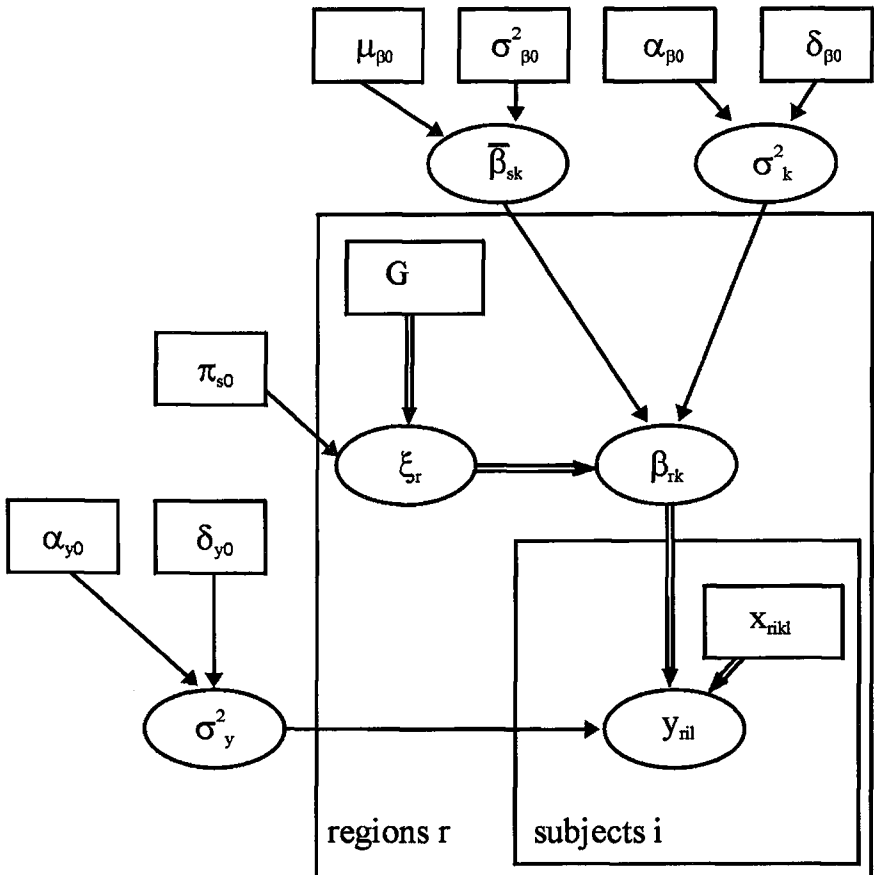
and

$$(5.8) \quad \sigma_k^{-2} \sim \text{Gamma}(\alpha_{\beta 0}, \delta_{\beta 0}).$$

Here α_{y0} , δ_{y0} , $\mu_{\beta 0}$, $\sigma_{\beta 0}^2$, $\alpha_{\beta 0}$, and $\delta_{\beta 0}$ are fixed hyper-parameters that will later be chosen so as to minimize the influence of prior information on the posterior distributions of the parameters. A schematic overview of the model, with the dependencies between the parameters is provided in Figure 5.1. The figure presents a *directed acyclic graph* (cf., Best et al. 1996), which illustrates the hierarchical structure of the model and shows the conditional relations required setting up a Gibbs estimation procedure.

Figure 5.1

Schematic overview of the geographic segmentation model with constraints



5.2.3 Model Estimation and Determination of Number of Segments

In estimating the geographic segmentation model, we take a Bayesian approach. The general objective of a Bayesian approach is to derive the posterior distribution of the parameters, given the data. Instead of estimating a single parameter value, a probability distribution of the parameters is obtained from the information contained in the data. In practice, however, this posterior distribution cannot be derived in closed form and is often unknown. This is in particular the case for the geographic segmentation model outlined above. To overcome this problem, we develop a Gibbs sampling scheme (Casella and George 1992; Gelfand and Smith 1990; Grenander 1983). The Gibbs sampler identifies the desired posterior distribution of the parameters, by sampling each parameter from its so-called full conditional distribution, i.e., the distribution of a single parameter (or set of parameters), given all other parameters and data. As opposed to the (unconditional) posterior distribution, these full conditional posterior distributions can be derived in closed form. Since the prior distributions of the parameters are taken from conjugate families, the posterior distributions can be derived and reduce to normal and gamma distributions, as follows:

$$(5.9) \quad \beta_r \sim N_K \left(\Sigma_*^{-1} \left(\Sigma^{-1} \bar{\beta}_s + \sigma_y^{-2} X_r' y_r \right) \Sigma_*^{-1} \right),$$

with vectors $\beta_r = (\beta_{rk})$, $\bar{\beta}_s = (\bar{\beta}_{sk})$, matrices $y_r = (y_{ril})$, $\Sigma = \text{diag}(\sigma_1^2, \dots, \sigma_K^2)$, $\Sigma_* = \sigma_y^{-2} X_r' X_r + \Sigma^{-1}$, and 3-way arrays $X_r = (x_{rikl})$,

$$(5.10) \quad \bar{\beta}_{sk} \sim N \left(\frac{\sigma_k^{-2} \sum_{\{r: \xi_r = s\}} \beta_{rk} + \sigma_{\beta 0}^{-2} \mu_{\beta 0}}{L_s \sigma_k^{-2} + \sigma_{\beta 0}^{-2}}, \frac{1}{L_s \sigma_k^{-2} + \sigma_{\beta 0}^{-2}} \right),$$

$$(5.11) \quad \sigma_k^{-2} \sim \text{Gamma} \left(\alpha_{\beta 0} + R/2, \delta_{\beta 0} + \sum_{r=1}^R (\beta_{rk} - \bar{\beta}_{\xi_r, k})^2 / 2 \right),$$

$$(5.12) \quad \sigma_y^{-2} \sim \text{Gamma}(\alpha_{y0} + N/2, \delta_{y0} + \sum_{ril} (y_{ril} - \sum_k x_{rik} \beta_{rk})^2 / 2),$$

$$(5.13) \quad p(\xi_r = s) = \begin{cases} \frac{\pi_{0s} p(\beta_r | \bar{\beta}_s, \Sigma)}{\sum_{s': s' \text{ adjacent to } s} \pi_{0s'} p(\beta_r | \bar{\beta}_{s'}, \Sigma)} : \Xi \in \Omega \\ 0 : \text{otherwise,} \end{cases}$$

$$\text{with } p(\beta_r | \bar{\beta}_s, \Sigma) = \prod_{k=1}^K p(\beta_{rk} | \bar{\beta}_{sk}, \sigma_k^2).$$

The Gibbs sampling scheme is defined in Equations (5.9) to (5.13). Using random starting values that comply with the restrictions on the segment memberships, the parameters are iteratively sampled from the full conditional distributions. In Equation (5.13) the segment-indicators ξ_r are sampled conditionally on the restrictions defined in Equation (5.5). We use a simple rejection procedure to take these restrictions into account (cf., Gelman et al. 1995). After sampling ξ_r , the restrictions on the segments are examined. If the segments are admissible, the sampled value of ξ_r is retained, otherwise new values of ξ_r are sampled until the restriction is satisfied. Note that in this way, we are able to estimate models with an extensive set of restrictions on the segments (see the discussion section of this chapter).

In estimating the model, we choose to run the Gibbs sampler for 15,000 iterations and discarded the first 3000 samples after visually inspecting whether the chains have converged. The remaining 12,000 samples were used to compute the sampling quantities of interest. We computed the median, and 5th and 95th percentiles of the posterior distributions of the parameters. In order for the priors not to influence the posterior estimates, we use diffuse, weakly-informative prior distributions in Equations (5.9) to (5.13): $\alpha_{y0} = \delta_{y0} = \alpha_{\beta0} = \delta_{\beta0} = 0.01$, $\sigma_{\beta0}^2 = 100$, and $\mu_{\beta0} = 0$.

Note that the model is defined conditionally on the number of segments in the market. Therefore, we estimate the model for an increasing number of segments and choose the appropriate number of segments using Bayes factors (Kass and Raftery 1995). The Bayes factor

is asymptotically equivalent to information criteria such as BIC and CAIC, that are frequently used in the marketing literature to decide on the appropriate number of segments.

5.3 Synthetic Data Analysis

In order to assess the performance of the model, we conducted a synthetic data analysis. There are three basic sources of error that may affect parameter recovery: (1) error caused by differences within segments (within-segment heterogeneity), (2) error due to unexplained variance of the dependent variable, and (3) error induced by not observing all subjects within regions (sampling error). Based on these factors, we generated data sets according to the following design:

- WVAR:** Within-segment heterogeneity, $\sigma_k = .10$ or $\sigma_k = .75$.
- EVAR:** Variance explained by the regressions, 90% or 50%, which corresponds to $\sigma_y^2 = .02$ and $\sigma_y^2 = .20$, respectively.¹
- RSIZE:** The number of subjects per region, fixed ($I_r = 42$) versus variable ($I_r \in [5, 80]$).

For each combination of the factors we generated a data set, which resulted in eight data sets to be analyzed. The total number of observation was equal to 4200 and the number of regions was 100. Two contiguous geographic segments were constructed by assigning half of the regions to the first segment ($\xi_r = 1$) and the remaining regions to the other segment ($\xi_r = 2$). Regressors $\{x_{rikl}\}$ were taken from the interval $[-.50, .50]$ and the $\bar{\beta}_{sk}$ from $[0, 1]$ ($k = 1, \dots, 6$), including an intercept. Then, region-specific parameters β_{rk} were drawn from normal distributions with means $\bar{\beta}_{\xi_r, k}$

¹ Note that there is a relation between the explained variance R^2 and the residual variance: $\sigma_y^2 \approx (\beta'X'X\beta / R^2 - \beta'X'X\beta)/(n-k)$.

and variance σ_k^2 , with values of $\bar{\beta}_{\xi_r, k}$ and σ_k^2 dictated by the particular experimental condition of the experimental design. Finally, for each subject a single dependent variable y_{ri1} was sampled ($L=1$) from the normal distribution with mean $\sum_{k=1}^K x_{rik1} \beta_{rk}$ and variance σ_y^2 .

We applied our model to each of the eight data sets. We iteratively sampled 3000 values from the conditional distributions in Equations (5.9) to (5.13). The first 1000 iterations were discarded and the median values of the parameters were calculated for the remaining 2000 values to reflect the posterior estimates. For each cell of the design we calculated a number of measures of model performance. We computed the hit rate of the segments as the percentage of correct predictions of the segment memberships ξ_r . For the parameters $\bar{\beta}_{sk}$, β_{rk} , σ_k^2 , and $\hat{\sigma}_y$, we computed the root mean squared error (RMSE) of the estimated versus the true parameters, taking the mean over the regressors k and regions r .

In Table 5.1 we report the performance measures of each of the eight data sets. The model seems to predict the segment-memberships very well. The hit rate exceeds 92 percent on average, which means that in more than 92 percent of the cases, the model assigns the regions to the correct geographic segment. The assignment of regions to segments deteriorates as the within-segment variance (*WVAR*) increases, especially with a variable number of observations per region (*RSIZE*). An increasing level of error variance (*EVAR*) does not seem to affect the hit rate. Recovery of the segment parameters, $\bar{\beta}_{sk}$, is only affected by the within-segment variance. This can be explained from the fact that the within-segment variance increases the overlap of segments. The $\text{RMSE}(\bar{\beta}_{sk})$ is very low (about .03) in case the within-segment variance is limited ($\text{WVAR}=.10$) and increases to about .13 if $\text{WVAR}=.75$. The recovery of $\bar{\beta}_{sk}$ is hardly affected by the variability of the sample sizes across regions and the explained variance of the regressions.

Table 5.1

*Model performance measures for the synthetic data sets**

FACTORS			PERFORMANCE MEASURE				
RSIZE	EVAR	WVAR	Hit rate	RMSE			
			ξ_r	$\bar{\beta}_{sk}$	β_{rk}	σ_k	σ_y^a
42	.02	.10	100 %	.0249	.0663	.0157	.0022
42	.02	.75	90 %	.1255	.0756	.0631	.0005
42	.20	.10	99 %	.0407	.1064	.0120	.0078
42	.20	.75	83 %	.1062	.0611	.0611	.0041
[5,80]	.02	.10	100 %	.0249	.1322	.0157	.0022
[5,80]	.02	.75	84 %	.1465	.9201	.1786	.0167
[5,80]	.20	.10	100 %	.0344	.1108	.0185	.0045
[5,80]	.20	.75	84 %	.1501	.7393	.1921	.0954

* All original parameter values were contained in the 95 percent Bayesian confidence intervals of the parameter estimates; ^a Note that we had a single parameter estimate for σ_y in each cell of the design. Therefore, the RMSE of this parameter reduces to the absolute error.

As expected, the precision of the region-specific parameters (β_{rk}) is affected by the variability of the sample sizes across regions. If the sample size is small in certain regions, the precision of the parameter estimates deteriorates. This effect increases with higher levels of within-segment variance. Apparently, the β_{rk} are shrunk towards their segment means $\bar{\beta}_{sk}$, which better represent their original values when differences within segments are small. This effect seems to carry over to the identification of σ_k , and σ_y . The RMSEs of those parameters are primarily affected by variable sample sizes, with higher levels of within-segment variance.

In general, the explained variance hardly affects the recovery of the model parameters. The RMSEs and hit rates of the regressions explaining about 50 percent of the variation in the dependent variable ($EVAR=.20$), are close to those explaining 90 percent ($EVAR=.02$). The within-segment variance and the variation of segment sizes among regions affects the precision of the parameter estimates. Still, the synthetic data analysis suggests that the model was very well capable of recovering the spatial segmentation structure and the average precision of the parameter estimates was high (but was of course lower for the region-specific parameters).

5.4 Empirical Application

5.4.1 Background of the Study

We illustrate the methodology with an international store-image segmentation study on meat outlets in Europe. In retailing, international expansion has become an increasingly important strategy to attain growth (Bell, Davies, and Howard 1997; Kumar 1997; Sternquist and Kacker 1994). Examples of retailers expanding their chain formats to foreign markets are Ikea, Toys'R'Us, Carrefour, and Wal-Mart. Such companies typically open up new outlets in new geographic areas where they communicate their distinct positioning messages. To become successful abroad, the first step is to identify a viable target area that meets a particular retail positioning strategy. Such areas may extend to parts of a country, but may also span (parts of) multiple countries. It is crucial though that a segment consists of a single contiguous area, if only for logistic reasons, so we impose the contiguity restrictions on the target areas we intend to identify.

In retailing, a relevant and distinctive positioning is frequently realized through the development of a particular store image (Nevin and Houston 1980; Samli 1989). Store image development includes strategies such as (everyday) low pricing (e.g., discount stores such as Wal-Mart), emphasis on a particular aspect of a store's inventory (wide or deep),

improved service, atmosphere and/or quality (James, Durand, and Dreves 1976; Mazursky and Jacoby 1986). Store image is an important aspect in the development of positioning strategies for retail chains. In addition, store image has served as an important basis for segmentation of retail customers (e.g., Steenkamp and Wedel 1991). Therefore, we use the importance that consumers attach to store image attributes as a basis for identifying geographic target segments. The setting of the application is in the food-retailing domain. With about \$750 billion in sales a year, the food retailing industry is among the largest industries in the European Union (M+M Eurodata 1996). The top-ten retailers are already present in most EU countries and international expansion is still increasing.

5.4.2 Data

The data collected for this study were part of a larger survey on consumer behavior with respect to meat in the European Union. Mail questionnaires were sent out to members of an international household panel in seven countries of the European Union. Store image measures were obtained according to the methodology of Steenkamp and Wedel (1991).² The store image attributes included in the survey are assortment, pricing, product quality, service quality, and store atmosphere. These attributes are widely accepted as being relevant to store image (Steenkamp and Wedel 1991). In addition, we considered distance as a factor influencing store preference. Perceptions on distance and image attributes of primary meat outlets were measured on single item bipolar scales (see Table 5.2). Overall evaluations of primary meat outlets were measured using a two-item bipolar scale. This instrument presupposes a multi-attribute model of store image formation, with overall evaluations of stores as a dependent variable and image perceptions and distance as regressors in Equation (5.1). The β_k reflect the importance of the image attributes and distance. For each respondent a single observation was obtained ($L=1$).

² Other measurement techniques have been proposed, such as conjoint (Louviere and Johnson 1991). These techniques assume, however, that all respondents are familiar with a similar set of retail chains. In Europe, however, pan-European retailing is still in its infancy, which renders such approaches infeasible.

Table 5.2

Items for perceptions on image attributes, distance, and overall store evaluations

<i>Attribute perceptions</i>		
of very low quality		of very high quality
very bad service		very good service
very pleasant atmosphere		very unpleasant ambiance
very little variety in meat		very much variety in meat
very expensive		very cheap
very far away		very close by
<i>Overall store evaluations</i>		
very negative		very positive
very bad		very good

Before the data were collected, extensive cross-national pretests were conducted. First, the store image instrument was tested among 42 Dutch consumers for wording, interpretability, and layout and appropriate adjustments were made. In a second stage, the questionnaires were refined in pretests conducted in France, the Netherlands, and Spain (N=99). The fieldwork was carried out in 1996 by a pan-European marketing research agency, GfK-Europanel. The countries included in the main study were Belgium, France, Germany, Italy, the Netherlands, Portugal, and Spain. Throughout the entire process, back-translation procedures were used to ensure that the content of the statements was similar across languages (Brislin 1970).

The total sample comprises about 2000 consumers in 120 pre-specified regions (see Table 5.3). These regions are defined according to the *Nomenclature des Unités Territoriales Statistique* classification at level 2 (NUTS2), which includes the *regierungsbezirke* in Germany, the *provincies* in the Netherlands and Belgium, the *régions* in France, the *comunidades autonomas* in Spain, the *comissaoes de coordenação regional* and *regioes autonomas* in Portugal, and the *regioni* in Italy (cf., Eurostat 1987).

Table 5.3
Sample characteristics

Country	Sample size	
	<i>Subjects</i>	Regions
Belgium	285	9
France	298	21
Germany	320	40
Italy	234	18
Netherlands	309	12
Portugal	255	5
Spain	265	15
Total	1966	120

In some regions the number of degrees of freedom are insufficient to perform separate regression analyses, which renders a sequential procedure infeasible. Our Bayesian approach, however, borrows information from the complete sample to identify the model parameters. A region-specific parameter is shrunk towards the mean of the segment to which it belongs and is therefore identifiable even if the number of observations in that region is limited (cf., Blattberg and George 1991).

5.4.3 Results

We applied the geographic segmentation model to the data, imposed contiguity restrictions in Equation (5.4) and constrained the segments to consist of at least 5 regions ($N_s \geq 5$). The Bayes factor was minimal for a 5-segment solution, hence we describe the results of this solution. In Table 5.4 we report the posterior medians and 90 percent Bayesian confidence intervals of the posterior distribution of $\bar{\beta}_{sk}$.³ For 25 out of

³ A Bayesian analogue of a frequentist confidence interval is usually referred to as a credible set. In line with the terminology used by Carlin and Louis (1996), we will use the term confidence interval. Note, however, that as opposed to traditional frequentist confidence intervals, Bayesian credence intervals have a more natural interpretation. The probability of the parameter falling in the

the 35 parameters, the confidence intervals were completely contained in the positive domain and the sign of these parameters were correct. Whereas the five median posterior values of the price parameters as well as the confidence intervals are small, four of them have the correct sign. The small values are consistent with previous studies that found weaker price-effects in Europe as compared to other continents (e.g., Rao and Steckel 1995; Tellis 1988).⁴

Table 5.4

Median posterior values and 90 percent confidence intervals of $\bar{\beta}_{sk}$

	<i>T1</i>	<i>T2</i>	<i>T3</i>	<i>T4</i>	<i>T5</i>
Assortment	.213 (.11, .32)	.219 (.14, .30)	.176 (.08, .25)	.250 (.18, .33)	.201 (.15, .25)
Pricing	.052 (-.02, .13)	-.004 (-.06, .05)	.071 (-.03, .14)	.012 (-.07, .09)	.018 (-.04, .09)
Product quality	.067 (-.06, .17)	.206 (.14, .27)	.207 (.12, .29)	.341 (.25, .44)	.126 (.07, .18)
Service quality	.265 (.14, .39)	.259 (.14, .37)	.169 (.04, .32)	.147 (.03, .29)	.304 (.23, .37)
Store atmosphere	.394 (.28, .51)	.205 (.12, .29)	.198 (.11, .28)	.128 (.00, .25)	.220 (.15, .30)
Distance	-.027 (-.10, .04)	.059 (.01, .11)	.044 (-.01, .09)	-.036 (-.10, .04)	.063 (.02, .10)
Intercept	.422 (-.25, 1.01)	.458 (.16, .76)	1.067 (.64, 1.64)	1.043 (.57, 1.54)	.598 (.21, .99)
Segment size ^a	.12 (.09, .13)	.15 (.10, .21)	.26 (.24, .32)	.18 (.17, .21)	.29 (.26, .34)

^a Proportion of regions belonging to each segment.

Bayesian interval is equal to .90, a statement that cannot be made in a frequentist setting.

⁴ Another potential explanation for the non-significant price effect is a limited amount of variation in the measured price perceptions. However, the variances of these perceptions were of the same magnitude as the other image perceptions.

In addition to the segment-specific parameter estimates, we graphically present the locations of the geographic segments in Figures 5.2 to 5.6. The color intensity assigned to a particular region is higher for increasing levels of the posterior probability of that region belonging to a segment, i.e., $P[\xi=s \mid S, \text{observed data}]$. As can be observed in the figures, the segments are not completely separated, but have fuzzy boundaries.

Figure 5.2

Spatial location of geographic segment T1

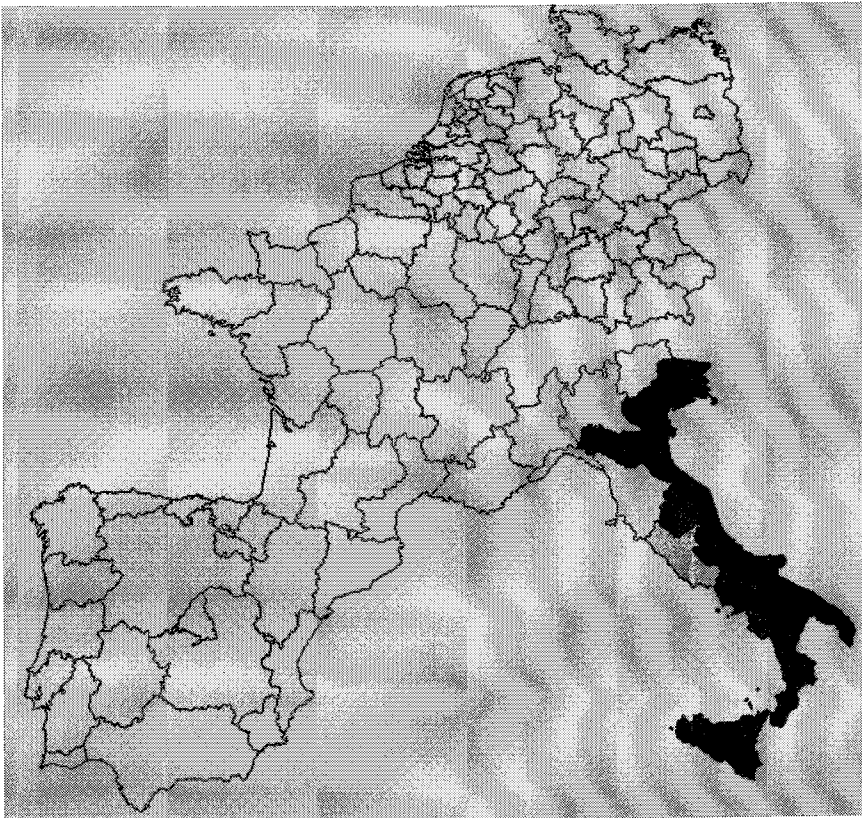


Figure 5.2 indicates that Segment *T1* is located within a single country, Italy, with high segment membership probabilities in the regions on the east coast. The northern part of Italy does not belong to this segment, which is consistent with the large cultural, economic, and

political differences that exist between the north and the south of Italy. With posterior distributions well in the positive domain, Segment *T1* exhibits large effects for service quality, store atmosphere and assortment (Table 5.4). The posterior distributions of the coefficients for price, product quality, and distance are concentrated around zero, which means that store positioning on those attributes will not be very effective in gaining appeal to consumers in this region. The most important attribute is store atmosphere, its median importance is close to .4. A retailer deciding to target this geographic segment may position its chain as having a pleasant store atmosphere, reconciling the positions taken by competitors in the area. A regional distribution center may be located in the central regions of Molise or Puglia, which provides the possibility to efficiently supply stores in the north and south of the geographic segment. A retailer may also use these central regions as an initial location of entry and gradually rollout towards the regions in the north and south of the segment.

Segment *T2* is a typical coastal area and transcends national borders (Figure 5.3). It includes Portugal, the western parts of Spain, and the Mediterranean coastal areas of Spain and France. Regions in central France, central Spain and west Italy have moderate segment memberships (well below .5). In this segment, strong similarities exist between regions in different countries, whereas the similarities between certain regions within countries are weak. For example, the coastal areas in the northeast of Spain are connected to the French Riviera but the French Riviera is not connected to the interior of France. This segment provides opportunities for cross-border retail operations, targeting an area that covers parts of different countries. Except for price, the image attribute estimates are all substantial. This segment bases its store image on service quality, but product quality, store atmosphere, and assortment cannot be neglected. The importance of distance of the stores to communities is significant, although the size of the associated coefficient is modest (.06).

Figure 5.3
Spatial location of geographic segment T2

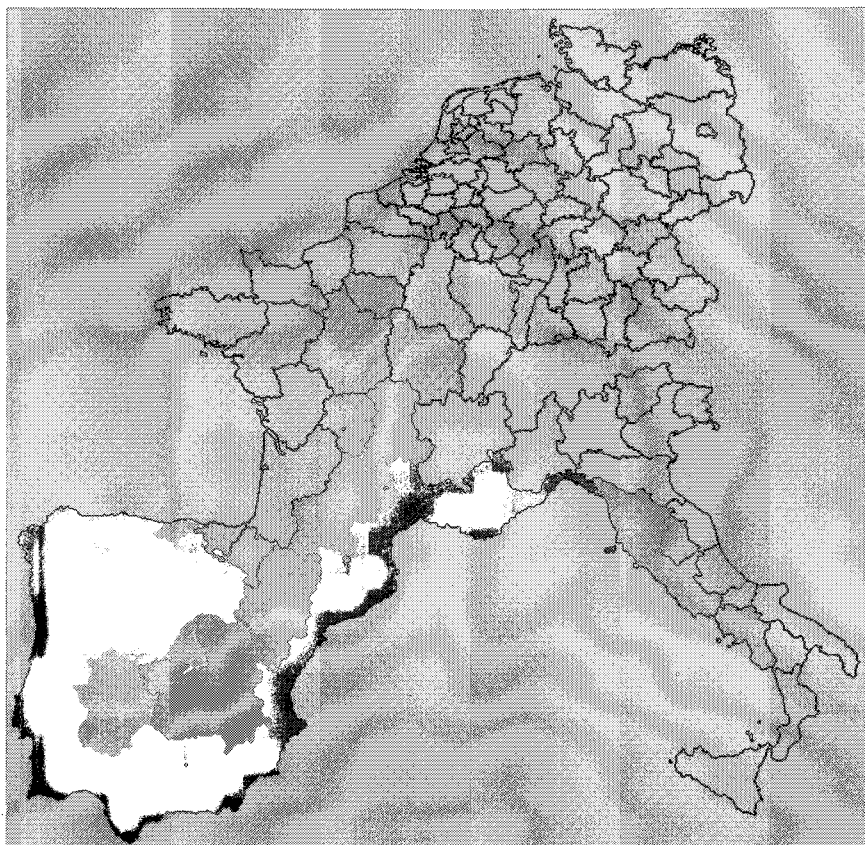


Figure 5.4
Spatial location of geographic segment T3

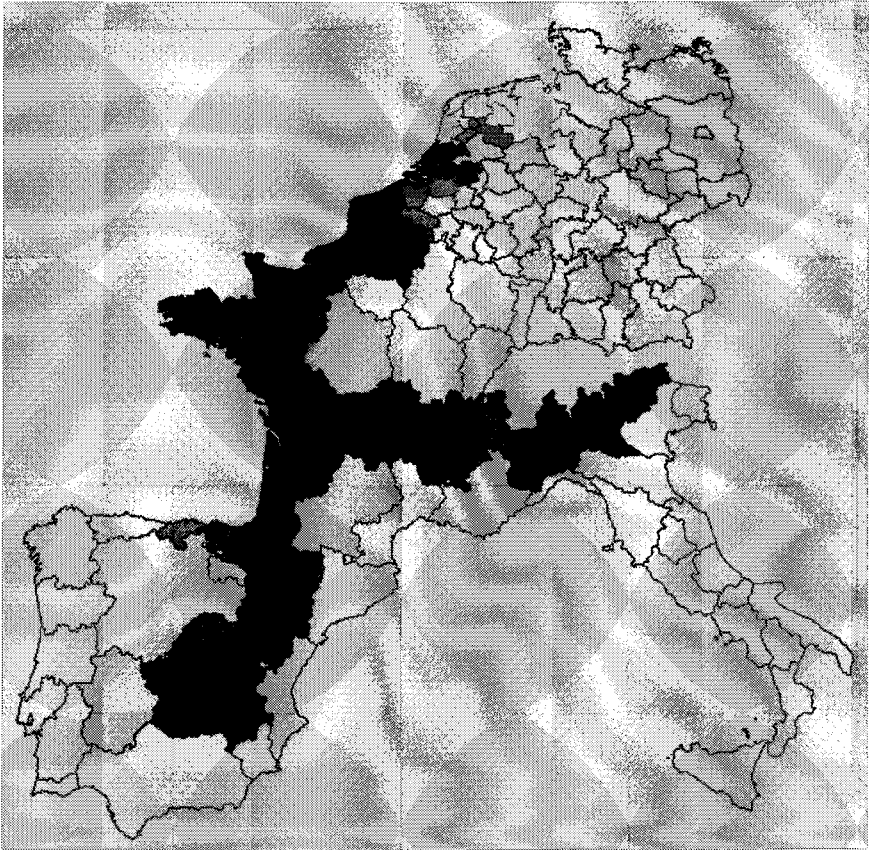
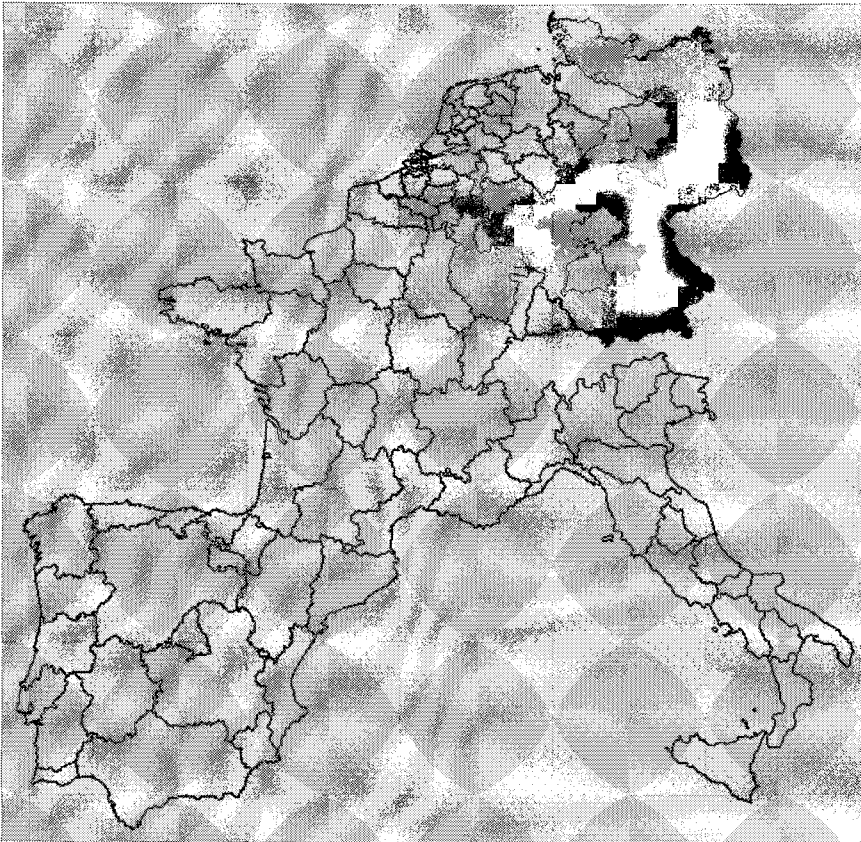


Figure 5.4 shows that segment *T3* covers parts of Belgium, France, Italy, the Netherlands, and Spain. The segment is mainly located at the west coast of France and stretches out towards central Spain, northern Italy, and the Netherlands. Interesting is again the clear border between northern Italy and its remainder. The northern part of Italy is united with the Lyon area, with which it historically shares strong economic and cultural ties. Segment *T3* is not characterized by one single important image attribute. Price and distance have negligible effect on store evaluations and are not important to consider in formulating expansion strategies. The posterior estimates of product quality, service quality, and

store atmosphere are much stronger and reliable (their confidence intervals are in the positive domain). Corresponding to segment T2, the image attribute importances are substantial but do not discriminate much.

Figure 5.5

Spatial location of geographic segment T4

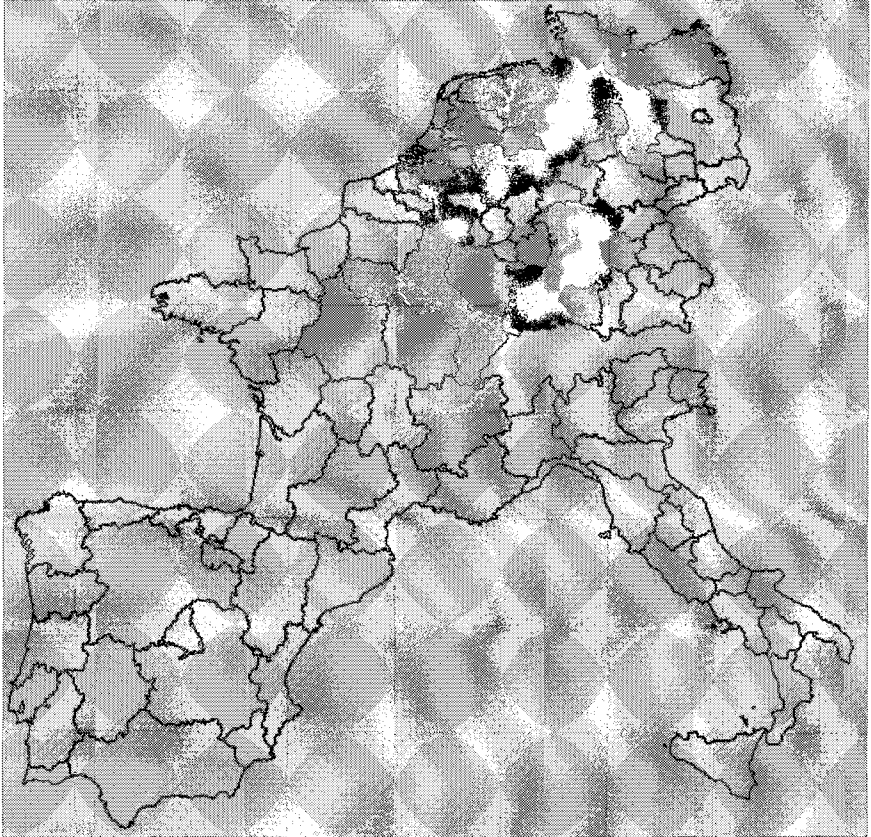


Segment T4 is mainly located in the former German Democratic Republic with offshoots to Bavaria in the south of Germany and the border regions where Belgium, Germany and France meet (Liege, Trier, and Lorraine; see Figure 5.5). This segment can efficiently be served by locating a distribution center in the central region of Thüringen. The most important

image attributes in this segment are product quality and assortment. This segment provides opportunities for upscale hypermarkets, such as Carrefour, that carry a wide assortment of high quality products.

Figure 5.6

Spatial location of geographic segment T5



The location of the final segment, *T5*, is depicted in Figure 5.6. In this segment store image is predominantly based on service quality and to a less extent on store atmosphere and assortment. The effect of distance is moderate but significant. The segment covers part of northwest Europe, including the Netherlands, northeast France, southwest and northwest Germany and parts of Belgium. The geographic configuration of this

segment lends itself for geographic staggered rollouts across borders. For example, a Dutch retail chain looking for expansion and which core competency is high on delivering service quality may expand into northwest Germany, initially penetrating markets in the regions of Detmold and Arnsberg, and gradually rollout towards Lüneburg, Hamburg, and Magdenburg. Similar approaches are possible to expand to the eastern regions of France (e.g., Champagne-Ardenne and Franche-Compte) through eastern Belgium. Still, the geographic configuration of this segment is less appealing than the other segments.

5.4.4 Segment Profiles

To further add to the managerial understanding of the geographic segmentation structure in terms of their differential accessibility, we construct additional profiles of the segments by relating the segments to secondary regional data. Secondary information on geodemographics and logistic accessibility were obtained from Eurostat, the central bureau of statistics of the European Union (Eurostat 1997). Additional data on media consumption were obtained from the international survey that contained the store image measurement instrument. The media consumption data provide an understanding of the relative accessibility of the segments through print, radio, and television advertising.

Because the location of the segments is not obtained with complete certainty (the boundaries of the segments are fuzzy), we construct profiles that accommodate the unreliability of the location of the geographic segments. For each partition generated by the Gibbs sampler, we computed the segment-specific averages of the variables. Then we computed the median and 95 percent coverage intervals across the iterations of the Gibbs sampler. The probability that the real values of the variables are contained in the coverage intervals is equal to .95. The results are given in Tables 5.5 and 5.6.

Table 5.5
*Regional characterization of the segments**

<i>Variables</i>	<i>T1</i>	<i>T2</i>	<i>T3</i>	<i>T4</i>	<i>T5</i>
<i><u>Geodemographics</u></i>					
Total area ^a	79.44 (59,82)	196.37 (127,271)	229.84 (192,288)	77.98 (68,95)	120.51 (93,147)
Population size ^b	39.46 (25,41)	45.82 (25, 65)	73.73 (65, 92)	33.86 (33, 45)	71.72 (56, 80)
Purchasing power ^c	14.83 (14, 15)	12.08 (10, 15)	16.25 (16, 18)	14.77 (14, 16)	18.39 (17, 19)
<i><u>Logistic accessibility</u></i>					
Motorway density ^d	13.27 (13,14)	7.46 (5,9)	10.65 (9,12)	18.90 (16,19)	18.04 (17,21)
Total traffic density ^e	59.38 (51,60)	44.13 (36,61)	30.89 (27,36)	86.22 (71,88)	56.43 (41,72)

* 90% coverage intervals in parentheses; ^a ×1,000 mile²; ^b ×1,000,000 inhabitants;

^c ×1,000 Euro per capita; ^d miles motorway per 1000 mile²; ^e number of vehicles per 1000 mile roads and highway.

Segment *T1* is relatively small and is characterized by an average pattern of geodemographics and logistic accessibility variables. This segment has a relatively high exposure to television broadcasts and is therefore well accessed through television advertising. Segment *T2* covers about 200,000 square miles of less developed area. The segment has the lowest population density and is less attractive in terms of purchasing power (20 percent below average) and logistic accessibility. Segment *T3* is the largest segment in terms of area and population. The limited traffic density may facilitate the logistic accessibility, but is cancelled out by the relatively low density of motorways. Segment *T4* covers the smallest area and has the least inhabitants. It is well developed in terms of motorways, but suffers from high traffic density. This segment may be well accessed through print advertising. Segment *T5* is an attractive segment to target with high population density and purchasing power. This segment is well accessible through radio advertisements and logistically well accessible because of its high density of motorways and modest traffic density.

Table 5.6

*Media profile of the consumers comprising the segments**

<i>Exposure to medium</i>	<i>T1</i>	<i>T2</i>	<i>T3</i>	<i>T4</i>	<i>T5</i>
Print advertising ^a	.99 (.9,1.0)	.72 (.7,.8)	.90 (.8,1.0)	1.39 (1.4,1.4)	1.09 (1.0,1.2)
Radio advertising ^b	9.9 (9,10)	11.0 (11,11)	14.8 (13,16)	16.0 (15,16)	16.7 (16,18)
TV advertising ^c	21.97 (22,23)	20.17 (20,21)	20.31 (20,21)	19.91 (19,20)	20.18 (20,21)

* 90% coverage intervals in parentheses; ^a Average number of daily newspapers; ^b number of hours per week of listening to the radio; ^c number of hours per week of watching TV.

5.4.5 Responsiveness of the Geographic Segments

Imposing contiguity restrictions on the geographic segments may jeopardize the responsiveness of the segments to a retail chain's image. Whereas the previous two sections have demonstrated the existence of differences between segments, we will now focus on the differences that exist within segments. The contiguity constraints may therefore yield segments that are more heterogeneous and would detract from their responsiveness. For that purpose, we re-estimated the model without contiguity constraints. We applied the Gibbs sampler using Equations (5.9) to (5.13), sampling from Equation (5.13) unconditional on the restrictions, i.e., leaving out the rejection step. By dropping these restrictions the number of feasible partitions increases, which will allow the within-segment variances to decrease. Therefore, we expect the heterogeneity parameters, σ_k , to be smaller for the unrestricted model. The extent to which the parameters increase is an important issue that needs to be assessed empirically.

Table 5.7

Posterior means and 90% confidence intervals of σ_k for restricted and unrestricted models

<i>Attribute</i>	<i>Restricted Segmentation</i>	<i>Unrestricted Segmentation</i>	<i>Percent Decrease</i>
Assortment	.0402 (.0304,.0540)	.0385 (.0301,.0515)	4.23%
Pricing	.0478 (.0350,.0650)	.0455 (.0338,.0615)	4.81%
Product quality	.0393 (.0296,.0523)	.0379 (.0291,.0499)	3.56%
Service quality	.0406 (.0305,.0534)	.0382 (.0293,.0505)	5.91%
Store atmosphere	.0398 (.0300,.0527)	.0379 (.0292,.0500)	4.77%
Distance	.0391 (.0391,.0505)	.0386 (.0297,.0506)	1.51%

Table 5.7 provides the estimates of the within-segment heterogeneity parameters σ_k for restricted and unrestricted segmentation solutions as well as the 90% confidence intervals. The values of σ_k for the restricted model are small, which indicates a high responsiveness of the segments. The within-segment heterogeneity of the restricted and the unrestricted models are very close. The percent decrease in heterogeneity ranges from only 1.51 percent for distance to 5.91 percent for service quality. This means that by imposing the contiguity restrictions, the segments become only slightly more heterogeneous. In addition, the differences of the confidence intervals between the restricted and unrestricted solutions are negligible. So, managers may benefit from the advantages of reduced costs of entry and accessibility of the segments, by sacrificing only a slight fraction of the responsiveness of the segments.

5.5 Discussion

In this study we proposed a new methodology to identify geographic target segments in international markets. The approach combines strengths from country segmentation schemes and traditional country selection strategies, by imposing geographic constraints on segments. The geographic concentration of such segments increases the accessibility, whereas the differences of the importance estimates positively relate to responsiveness. The model captures differences that are likely to exist within geographic segments. We illustrated the methodology in the setting of international expansion of retail chains in seven countries of the European Union. Based on the importance of different attributes of store image, we identified five geographic segments. The segments were different both in terms of their location and their pattern of store image attribute importances. One geographic segment comprised only part of a single country and the other segments crossed national borders. This suggests that, as opposed to targeting countries separately, alternate geographic configurations may be identified that are more responsive and do not sacrifice the accessibility of the segments.

The segments were related to secondary data on geodemographics, logistic accessibility, and media consumption, which added to the managerial relevance of the geographic segments identified. In addition, we tested the model against an unrestricted model that does not take the contiguity of segments into account. The results showed that the restrictions did not significantly affect the within-segment heterogeneity, which suggests that the accessibility of the contiguous segments is not removed by increased differences between consumer needs. A synthetic data analysis indicated that the model is very well capable of recovering the geographic segments and the model parameters even with a substantial amount of within-segment heterogeneity, a limited number of observations within regions and little variance explained.

The methodology proposed in this chapter is very flexible and accommodates many applications in response-based segmentation, including segmentation studies of metric conjoint analysis, image research, and customer satisfaction. The methodology is easily extended

to other applications, such as choice-based segmentation. The method generalizes to *generalized linear models* (McCullough and Nelder 1989), which accommodate other distributions of the dependent variables, including the (multinomial) logit, probit, Poisson, gamma, and negative binomial distributions. Suppose dependent observations y_{ril} are distributed according to a member of the exponential family, $f(y_{ril}|\theta_r, \sigma_y^2)$.

Then Equations (5.1) and (5.2) are replaced by:

$$(5.14) \quad y_{ril} \sim f(y_{ril} | \theta_r, \sigma_y^2),$$

$$(5.15) \quad h(\theta_r) = \sum_k x_{rik} \beta_{rk} ,$$

with σ_y^2 the dispersion parameter and $h(\cdot)$ an appropriate link function, connecting θ_r to a linear combination of the regressors and parameters. In order to estimate this model, a similar Gibbs sampling scheme can be derived, as defined in section 5.2.3. Because of the conditional independence structure of the model, only the full conditional posterior distributions of β_{rk} and σ_y^2 will change and will depend on the choice of distribution function $f(y_{ril}|\theta_r, \sigma_y^2)$. This generalized formulation nests extensions to other measurement instruments and segmentation bases including value systems (Kamakura and Mazon 1991), conjoint choice (DeSarbo et al. 1992), or brand choice (Kamakura and Russel 1989).

Given the simplicity of rejection sampling in Gibbs estimation, it is possible to cover a broad set of restrictions. For example, trade restrictions between countries may be imposed to accommodate situations where national borders are closed, such as those between North and South Korea. Other relevant constraints may be formulated to assure the accessibility of segments with media and distribution networks. Secondary information is often available that helps accommodating particular strategic desirabilities of the segments. For example, segments need to be sufficiently large and should be economically viable (Douglas and Wind 1987). Suppose T additional restrictions ($t=1, \dots, T$) are defined on the geographic segments, based on Q additional variables given for

each of the R regions. The additional variables are contained in an $(R \times Q)$ -matrix Z , which will typically include measures on sales potential per region such as household income, and purchase power. Let restriction t be satisfied if $A_t(Z, \Xi) \geq 0$. For example, let z_r the sales potential in region r then $A_t(Z, \Xi) \equiv \sum_{\{r: \xi_r = s\}} z_r - c$ ensures that segment s will be characterized by a sales potential equal to or exceeding c . The set of partitions that comply with these restrictions is defined as follows:

$$(5.16) \quad \Omega' = \{\Xi \mid A_t(Z, \Xi) \geq 0, \forall t = 1, \dots, T\}.$$

The set of feasible partitions that comply with both contiguity and other type of restrictions is contained in $\Omega \cap \Omega'$, with Ω as defined in Equation (5.4). In model estimation Ω in Equation (5.13) is replaced by $\Omega \cap \Omega'$. Finally, the shape of the segments can be controlled by making slight alterations to the contiguity restriction. Instead of taking the sum in Equation (5.4) over $N_s - 1$ regions, a smaller number can be chosen, e.g., $(N_s - 1)/2$. This restriction will lead to the identification of geographic segments with more concentrated shapes.

The empirical application was based on a large representative cross-national sample of almost 2000 consumers. The study covered a large area, extending to 120 regions in seven European countries. Additional care was taken to satisfy high quality requirements of the data. Several pretests were conducted and back-translation procedures were used to ensure a similar content of the statements across languages. Still, the study is not without its limitations. First, the study was tailored to outlets selling meat, such as butchers, and to meat inventory of convenience stores, and meat departments in supermarkets and hypermarkets. Second, the configuration of the countries included in the study was also limited. Italy was connected with the other countries through only two regions. The lack of available data in Switzerland and Austria may have obstructed the formation of segments covering north Italy and south Germany. In addition, the countries were part of a continent (Europe) that has a vast cultural history. Other "younger" continents such as the United States may display less spatial patterns in consumption. Third, criticism may be levied against the measurement

instrument used to derive the importances of image attributes. Perceptions of image attributes and distance of the stores were measured using single item measures. The use of multi-item scales help to control measurement error of the perceptions and to correct for cross-national response bias (cf., Steenkamp and Baumgartner 1998). Still, the purpose of the empirical application was to illustrate the methodology. Future research may shed more light on these issues by applying the model in other settings, using different outlets, other countries, and improved measurement instruments.

Chapter 6

Discussion

The primary objective of this thesis is to develop and validate new methodologies to improve the effectiveness of international segmentation strategies. The current status of international market segmentation research reviewed in the introductory chapter provided a number of methodological and substantive issues that needed further attention. In the previous chapters, these issues were critically assessed and methodologies were developed as potential solutions. In this chapter the main conclusions are summarized and issues for further research are provided.

6.1 Summary and Conclusions

6.1.1 Two Directions in International Segmentation Research

In chapter 1, previous research in international segmentation was classified according to the three dimensions depicted in Figure 1.2. In the figure, the first dimension relates to the segmentation basis, the second to segmentation objects, and the third to segmentation methodology. All

three dimensions affect the effectiveness of international segmentation strategies. Two key research directions for improving the effectiveness of international segmentation were formulated along these dimensions.

The first direction concerns the integration of targeted product and communication strategies by linking product-specific bases with general consumer-level bases. In this thesis a new methodology is developed to identify cross-national market segments using means-end chain theory. Based on theory founded in consumer behavior, the means-end chain links values (a general consumer-level basis) with benefits and attributes (product-specific bases). Such an approach has the potential to combine product development and communication strategies at the international segment level and may serve as a guiding principle for international marketers to tailor products and advertising messages to the desires of global consumer segments. Chapter 4 provides a model-based methodology for identifying such segments. An international segmentation model was developed that estimates relations between product attributes, benefits of product use, and consumer values at the international segment level, and at the same time identifies those segments. The model builds upon methodological issues that were addressed in chapters 2 and 3 and rests on mixture methodology that, due to its capability of deriving segments based on models of consumer behavior, is particularly effective. In particular, it accounts for the international sampling design and the different response tendencies across countries and consumers.

The segmentation model was applied to identify segments in the European yogurt market, using a large sample of European consumers. Four segments were identified, of which one was truly pan-European and the other segments were cross-national. The segments were found to represent distinctive means-end structures and the pattern of links between attributes and values gave rise to strategic implications with respect to product development and communication. The segments were found to be related to socio-demographics, consumption patterns, media consumption, and personality data, which contributes to the identifiability and accessibility of the segments. The results suggest that the proposed model-based international segmentation methodology, combining

product- and consumer-level bases, has the potential to identify segments of consumers in different countries that are actionable towards product development and advertising strategy.

In chapter 5, a different direction is proposed that seeks to improve the effectiveness of international target market selection of expanding companies, by improving the geographic configuration of segments. Whereas consumer segments are more responsive, their typical geographic configuration does not make them accessible with cost efficient logistic operations. Especially if physical distribution represents a major component of total product costs, it is important that a geographic segment defines one particular area as opposed to dispersed segments that may arise in previous segmentation approaches. A flexible model-based segmentation approach is developed that identifies contiguous geographic segments based on consumer-level data. The model is based on multi-attribute theory of preference formation and accommodates a broad set of strategic restrictions on the segments. Moreover, the model accounts for heterogeneity that is likely to exist within geographic segments.

The methodology is illustrated in the international retailing domain, where geographic expansion is an important strategy to attain growth. Based on the importance that consumers attach to different attributes of store image, five geographic segments were identified across regions in seven countries of the European Union. The segments were distinctive in terms of their patterns of image attribute importances, which provides opportunities for expanding retailers to delineate geographic areas to enter and to develop an appropriate image in such areas. The results also demonstrated the accessibility of the segments through advertising media and logistics. In addition, no significant differences were found between the original model and a nested model that does not take the contiguity into account. This means that the actionability of restricting segments to be contiguous does not substantially harm the responsiveness of these segments.

6.1.2 Methodological Issues in International Market Segmentation

Given the often limited rigor of statistical and measurement techniques applied in the area of international segmentation, special attention has been given to methodological issues. Several issues were addressed that may negatively affect international segmentation research findings and methods were developed to deal with these issues.

International segmentation models. The first issue concerns the segmentation method. International segmentation research demonstrates an excessive reliance on heuristic segmentation techniques, such as cluster analysis. These techniques provide limited flexibility for international segmentation and may not be very effective in recovering response-based segments. The international segmentation methodologies developed in this thesis are model based and rely on insights from state of the art statistical techniques such as mixture and hierarchical Bayes models. Three international segmentation models are described in chapters 2, 4, and 5, and are successfully applied to empirical data. Chapter 5 provided a Bayesian formulation of a new international segmentation model that accommodates within-segment heterogeneity and complex restrictions on the configuration of segments. In chapter 4 it is empirically shown that a new mixture model approach outperforms standard clustering approaches that are traditionally employed in international segmentation.

Sampling designs. A second methodological issue is related to the estimation of international segmentation models. The importance of international sampling designs had not been acknowledged in the literature on international segmentation and mixture modeling. Previous international segmentation studies did not account for the implicit stratified sampling designs encountered in cross-national data collection. In this thesis the effects of international sampling designs on maximum likelihood estimation of segmentation models are investigated and a framework for accommodating those effects is proposed. A pseudo maximum likelihood procedure is introduced that accommodates

complex sample designs for maximum likelihood estimation of finite mixture models. In addition, modified or pseudo-information criteria are suggested for correct estimation of the number of international segments. The effects of not accounting for the sampling design were empirically assessed in an international value segmentation study. The pseudo-maximum likelihood approach was compared to standard maximum likelihood estimation that does not account for the sampling design. The results show that the estimates of segment sizes and segment-level parameters may be severely biased when not accounting for the design in standard maximum likelihood estimation. In addition, the empirical application demonstrated that the use of standard information criteria leads to incorrect inferences about the number of segments. This means that standard estimation methods in international segmentation research may lead to incorrect conclusions and erroneous managerial action.

Measurement. The international segmentation methodology in chapter 4 was based on MEC theory. The traditional method for measuring means-end chains (laddering) is not suitable for international segmentation. A necessary condition for the validity of international segments is that the basis for segmentation is measured in a valid and reliable way. Measurement instruments should allow collecting large and representative samples and standardization across countries. In this thesis a MEC measurement technique is developed that meets those criteria. The technique is denoted as the association pattern technique (APT), and its validity is further assessed.

Two key issues were investigated that may hamper the validity of APT. First, APT implicitly assumes that attribute-benefit and benefit-value links are independent because it measures these links in two separate tasks. The second issue is the convergent validity of APT as compared to the more traditional laddering interview. Consistent support for independence of attribute-benefit and benefit-value links was found across four product categories. Statistical tests of convergent validity of APT and laddering demonstrated that the basic structure revealed by both methods is similar. This suggests that APT is valid for measuring means-end chains and can be used for identifying international consumer

segments. The APT method is successfully applied in an international segmentation study in 11 countries.

Response tendencies. Response tendencies may hamper the identification of cross-national segments. The APT method may be prone to a respondent's propensity to choose any link. Therefore, the international segmentation model in chapter 4 accounted for differences in those tendencies that may exist between respondents. Based on item response theory, a response threshold approach was developed that allows testing those differences between countries, but also within countries. The results demonstrated that the differences in response tendencies were significant between countries, but also within countries. This means that it is important to account for response tendencies in international segmentation but in domestic segmentation as well.

6.2 Limitations and Issues for Future Research

In chapter 2, a method is presented that accommodates sampling designs in international segmentation models. The method is successfully applied in two international segmentation studies. A generalized PML method is presented, which accommodates various types of probability samples. The PML approach, however, is not without its limitations. PML provides an alternative for maximum likelihood-based segmentation models but is not directly applicable in a Bayesian framework. The Bayesian segmentation methodology in chapter 5 does not explicitly accommodate international sampling designs. However, the results are not affected by the design if the geographic regions that are partitioned into segments constitute the stratification units. Such designs are 'ignorable' because of the hierarchical structure of the model (cf., Gelman et al. 1995). For example, the model structure will level out the effect of the design if samples are stratified by county and the states are object of segmentation. Nevertheless, sampling designs different from the design described above may lead to incorrect inferences. Future research is needed to assess whether and to what extent the results of Bayesian

segmentation models are affected by sampling designs and how probability samples are accommodated in such models.

The segmentation model developed in chapter 4 requires the estimation of the segment-specific importances of each link. Because of the large set of parameters in the model, the Newton-type method that is employed for model estimation may be inefficient and may render the estimation of the asymptotic variance-covariance matrix unreliable. More efficient estimation approaches can be taken by estimating the model using MCMC methods as implemented in chapter 5. Such an approach is likely to improve the reliability estimates of the parameters and to speed up computation.

The international segmentation methods and bases developed in this thesis resulted in the recovery of actionable, accessible, and identifiable segments. It is important though to assess the real economic benefits of these approaches. The identification of international segments is only the first phase in the process of international target marketing. International companies need to choose which market segments to serve, based on the segment sizes, growth potential, and their attractiveness in the long run. This choice highly depends on a company's objectives and resources and on the fierceness of competition in these markets. In addition, companies need to develop differentiation strategies to position their products for the selected international market segments. A critical question is whether and under what conditions targeted international marketing strategies result in increased sales and profits, as compared to undifferentiated standardized strategies and other methods and bases for segmentation. Some research has been done in this area. Using the MEC segmentation results of chapter 4, new product concepts and marketing strategies were developed and validated with companies and consumers. The results were promising and suggested that international marketing strategies developed on the basis of the MEC approach to segmentation lead to viable and useful outcomes (see Steenkamp et al. 1998 for more details).

The empirical research in this thesis supports the existence of cross-national segments in the European Union. Substantial segments were identified, consisting of tens of millions of households from different European countries sharing similar needs. These segments provide

opportunities for international companies to tailor their products and other elements of their marketing mixes to large international segments instead of individual countries, benefiting from cost reductions through scale economies in production and marketing activities. A logical extension is to examine the existence of worldwide segments. Future research may focus on this issue and apply similar model-based approaches to countries from different continents and bring an end to the long lasting debate between proponents of Levitts' (1983) conception of worldwide customers and those adhering local adaptation.

An interesting development in the environment of today's marketers is the rise of the Internet, which provides a new way of selling and advertising to consumers worldwide (Hoffman and Novak 1996). Online retailers sell products through virtual stores and manufacturers sell and advertise their products directly to end customers. At this very moment, the population of Internet users is expanding to broader sections of the world population and the volume of online sales is growing at a breathtaking rate. The Internet is particularly important for international segmentation because it is a truly global medium. It brings together consumers and sellers from all over the world and has created the potential for companies of all sizes to target various global segments, ranging from individual consumers and small global niches to large worldwide segments. This requires marketers to develop new approaches to international segmentation of consumers on the Internet. An interesting example is provided by SRI Consulting, creator of the VALS 2 lifestyle typology, who already developed ten iVALS segment profiles of worldwide customers on the Internet (SRI 1999). Future research may shed more light on the effectiveness of existing approaches to segmenting online customers and may seek to develop new methods and bases for worldwide segmentation in cyberspace.

References

- Abbott, L. (1955), *Quality and Competition*. Westwood: Greenwood Press.
- Agresti, Alan (1990), *Categorical Data Analysis*. New York: John Wiley & Sons.
- Aitkin, Murray and Donald B. Rubin (1985), "Estimation and Hypothesis Testing in Finite Mixture Models," *Journal of the Royal Statistical Society B*, 47, 67-75.
- Akaike, H. (1974), "A new Look at the Statistical Identification Model," *IEEE Transactions on Automatic Control*, 19, 716-723.
- Alba, Joseph W. and J. Wesley Hutchinson (1987), "Dimensions of Consumer Expertise," *Journal of Consumer Research*, 13, 411-54.
- Alden, Dana L., Jan-Benedict E.M. Steenkamp and Rajeev Batra (1999), "Brand Positioning Through Advertising in Asia, North America and Europe: The Role of Global Consumer Culture," *Journal of Marketing*, 63 (January), 75-87.
- Allenby, Greg M., Neeraj Arora, and James L. Ginter (1998), "On the Heterogeneity of Demand," *Journal of Marketing Research*, 35 (August), 384-9.
- Allenby, Greg M. and James L. Ginter (1995), "Using Extremes to Design Products and Segments Markets," *Journal of Marketing Research*, 32 (November), 392-403.
- Allenby, Greg M. and Peter E. Rossi (1999), "Marketing Models of Heterogeneity," *Journal of Econometrics*, 89, 57-78.
- Amemiya, Takeshi (1985), *Advanced Econometrics*. Cambridge: Harvard University Press.

- Aptech (1995), *Constrained Maximum Likelihood, GAUSS Manual*. Maple Valley: Aptech Systems.
- Askegaard, Soren and Tage Koed Madsen (1998), "The Local and the Global: Exploring Traits of Homogeneity and Heterogeneity in European Food Cultures," *International Business Review*, 7 (6), 549-68.
- Aurifeille, Jaques-Marie and Pierre Valette-Florence (1995), "Determination of the Dominant Means-end Chains: A Constrained Clustering Approach," *International Journal of Research in Marketing*, 12 (3), 267-78.
- Baker, Frank B. (1992), *Item Response Theory*. New York: Marcel Dekker.
- Bass, Frank M. (1969), "A New Product Growth for Model Consumer Durables," *Management Science*, 15 (January), 215-27.
- Baumgartner, Hans and Jan-Benedict E.M. Steenkamp (1996), "Exploratory Consumer Buying Behavior: Conceptualization and Measurement," *International Journal of Research in Marketing* 13 (2), 121-37.
- Baumgartner, Hans and Jan-Benedict E.M. Steenkamp (1999), "Response Styles in Marketing Research: A Cross-National Investigation," working paper, Penn State University.
- Bell, Richard, Ross Davies, and Elizabeth Howard (1997), "The Changing Structure of Food Retailing in Europe: The Implications for Strategy," *Long Range Planning*, 30 (6), 853-61.
- Best, N.G., D.J. Spiegelhalter, A. Thomas, and C.E.G. Brayne (1996), "Bayesian Analysis of Realistically Complex Models," *Journal of the Royal Statistical Society A*, 159, 323-42.
- Bettman, James (1974), "Toward a Statistics for Consumer Decision Net Models," *Journal of Consumer Research*, 1, 71-80.
- Binder, David A. (1983), "On the Variances of Asymptotically Normal Estimators from Complex Surveys," *International Statistical Review*, 51, 279-92.
- Blattberg, Robert C. and Edward I. George (1991), "Shrinkage Estimation of Price and Promotional Elasticities: Seemingly Unrelated Equations," *Journal of the American Statistical Association*, 86 (June), 304-15.
- Bock, R. Darrell and Murray Lieberman (1970), "Fitting a Response Model for n Dichotomously Scored Items," *Psychometrika*, 35 (June), 179-97.
- Bonett, Douglas G. and Peter M. Bentler (1983), "Goodness-of-fit Procedures for the Evaluation and Selection of Log-linear Models," *Psychological Bulletin*, 93 (1), 149-66.

- Boote, Alfred S. (1983), "Psychographic Segmentation in Europe... A Case for Standardized International Advertising," *Journal of Advertising Research*, 22 (December/January), 19-25.
- Bozdogan, Hamparsum (1987), "Model Selection and Akaike's Information Criterion (AIC): The General Theory and its Analytical Extensions," *Psychometrika*, 52, 345-70.
- Bozdogan, Hamparsum (1990), "On the Information-Based Measure of Covariance Complexity and its Application to the Evaluation of Multivariate Linear Models," *Communications in Statistics, Theory and Methods*, 19, 221-78.
- Brislin, Richard W. (1970), "Back-Translation for Cross-Cultural Research," *Journal of Cross-Cultural Psychology*, 1, 185-216.
- Carlin, Bradley P. and Thomas A. Louis (1996), *Bayes and Empirical Bayes Methods for Data Analysis*, New York: Chapman & Hall.
- Casella, George and Edward I. George (1992), "Explaining the Gibbs Sampler," *The American Statistician*, 46 (August), 167-74.
- Cattin, Philippe, Alan E. Gelfand, and Jeffrey Danes (1983), "A Simple Bayesian Procedure for Estimation in a Conjoint Model," *Journal of Marketing Research*, 20 (February), 29-35.
- Chen, Chuansheng, Shin-ying Lee, and Harold W. Stevenson (1995), "Response Style and Cross-Cultural Comparisons of Rating Scales Among East Asian and North American Students," *Psychological Science*, 6 (May), 170-5.
- Claeys, C., A. Swinnen and P. Vanden Abeele (1995), "Consumers' Means-end Chains for "Think" and "Feel" Products," *International Journal of Research in Marketing*, 12 (3), 193-208.
- Clare, Donald A., and Donald G. Sanford, (1979), "Mapping Personal Value Space: A Study of Managers in Four Organizations," *Human Relations*, 32, 659-66.
- Cochran, W.G. (1977), *Sampling Procedures*, New York: John Wiley & Sons.
- Cooper, Robert G. (1984), "New Product Strategies: What Distinguishes the Top Performers?" *Journal of Product Innovation Management*, 2 (September), 151-64.
- Day, Ellen, Richard J. Fox, and Sandra M. Huszagh (1988), "Segmenting the Global Market for Industrial Goods: Issues and Implications," *International Marketing Review*, 5 (3), 14-27.

- Dempster, A.P., N.M. Laird, and R.B. Rubin, (1977), "Maximum Likelihood from Incomplete Data via the EM-algorithm," *Journal of the Royal Statistical Society series B*, 39, 1-38.
- DeSarbo, Wayne S. and W.L. Cron (1988), "A Maximum Likelihood Methodology for Clusterwise Linear Regression," *Journal of Classification*, 5, 249-82.
- DeSarbo, Wayne S. and Elizabeth A. Edwards (1996), "Typologies of Compulsive Buying Behavior: A Constrained Clusterwise Regression Approach," *Journal of Consumer Psychology*, 5 (3), 231-62.
- DeSarbo, Wayne S. and David K. Hildebrand (1980), "A Marketer's Guide to Log-linear Models for Qualitative Data Analysis," *Journal of Marketing*, 44 (Summer), 40-51.
- DeSarbo, Wayne S. and Vijay Mahajan (1984), "Constrained Classification: The Use of a Priori Information in Cluster Analysis," *Psychometrika*, 49 (June), 187-215.
- Diebolt, J. and Christian P. Robert (1994), "Estimation of Finite Mixture Distributions Through Bayesian Sampling," *Journal of the Royal Statistical Society B*, 56, 163-75.
- Dillon, William R., Donald G. Frederick and Vanchai Tangpanichdee (1985), "Decision Issues in Building Perceptual Product Spaces with Multi-attribute Rating Data," *Journal of Consumer Research*, 12 (June), 47-63.
- Dillon, William R. and Matthew Goldstein (1984), *Multivariate Analysis Methods and Applications*, New York: John Wiley & Sons.
- Douglas, Susan P., and C. Samuel Craig, (1992), "Advances in International Marketing: A Review," *International Journal of Research in Marketing*, 9 (4), 291-318.
- Douglas, Susan P. and C. Samuel Craig (1999), *International Marketing Research*. 2nd ed., forthcoming.
- Douglas, Susan P. and Christine Urban (1977), "Life-style Analysis to Profile Woman in International Markets," *Journal of Marketing*, 41 (July), 46-54.
- Douglas, Susan P. and Yoram Wind (1987), "The Myth of Globalization," *Columbia Journal of World Business* (Winter), 19-29.
- Eshgi, Abdolreza and Sheth, Jagdish N. (1985), "The Globalization of Consumption Patterns: An Empirical Investigation," in *Global Perspectives in Marketing*, Kaynak, Erdener, ed. New York: Praeger, 133-48.
- Eurostat (1987), *Regions: Statistical Yearbook*. Luxembourg: Office for Official Publications of the European Communities.

- Eurostat (1997), *Regio Database*. Luxembourg: Office for Official Publications of the European Communities.
- Fahrmeier, Ludwig and Gerhard Tutz (1991), *Multivariate Statistical Modelling Based on Generalized Linear Models*. New York: Springer.
- Farley, John U. and Donald R. Lehmann (1994), "Cross-National "laws" and Differences in Market Response," *Management Science*, 40 (1), 111-22.
- Feather, N.T. (1984), "Protestant Ethic, Conservatism, and Values," *Journal of Personality and Social Psychology*, 46, 1132-41.
- Ferligoj, Anuska and Vladimir Batagelj (1982), "Clustering with Relational Constraints," *Psychometrika*, 47 (December), 413-26.
- Ferligoj, Anuska and Vladimir Batagelj (1983), "Some Types of Clustering with Relational Constraints," *Psychometrika*, 48 (December), 541-52.
- Fienberg, S. (1980), *The Analysis of Cross-classified Categorical Data*. Cambridge: MIT Press.
- Frank, Ronald E., Massy, William F., and Wind, Yoram (1972), *Market Segmentation*. Englewood Cliffs: Prentice-Hall.
- Gelfand, Alan E. and Adrian F.M. Smith (1990), "Sampling-Based Approaches to Calculating Marginal Densities," *Journal of the American Statistical Association*, 85 (June), 398-409.
- Gelman, Andrew, John B. Carlin, Hal S. Stern, and Donald B. Rubin (1995), *Bayesian Data Analysis*. London: Chapman & Hall.
- Geman, Stuart and Donald Geman (1984), "Stochastic Relaxation, Gibbs Distributions, and the Bayesian Restoration of Images," *IEEE Transactions on pattern analysis and machine intelligence*, PAMI-6 (November), 721-41.
- Ghose, Sanjay (1998), "Distance Representations of Consumer Perceptions: Evaluating Appropriateness by Using Diagnostics," *Journal of Marketing Research*, 35 (May), 137-53.
- Ghosh, Malay (1995), "Inconsistent Maximum Likelihood Estimators for the Rasch Model," *Statistics and Probability Letters*, 23, 165-70.
- Goodman, L.A. (1970), "The Multivariate Analysis of Qualitative Data: Interactions Among Multiple Classifications," *Journal of the American Statistical Association*, 65, 226-56.
- Gordon, A.D. (1973), "Classification in the Presence of Constraints," *Biometrics*, 29 (December), 821-7.

- Green, Paul E. and V. Srinivasan (1978), "Conjoint Analysis in Consumer Research: Issues and Outlook," *Journal of Consumer Research*, 5 (September), 103-23.
- Green, Paul E., Frank J. Carmone, and David P. Wachspress (1976), "Consumer Segmentation via Latent Class Analysis," *Journal of Consumer Research*, 3 (December), 170-4.
- Grenander, U. (1983), *Tutorial in Pattern Theory*. Division of Applied Mathematics, Brown University.
- Griffin, Abbie (1992), "Evaluating QFD's Use in U.S. Firms as a Process for Developing Products," *Journal of Product Innovation Management*, 9, 171-87.
- Grunert, Klaus G. (1997), "What's in a Steak? A Cross-cultural Study on the Quality Perception of Beef," *Food Quality and Preference*, 8, 157-71.
- Grunert, Klaus G. and Suzanne C. Grunert (1995), "Measuring Subjective Meaning Structures by the Laddering Method: Theoretical Considerations and Methodological Problems," *International Journal of Research in Marketing*, 12 (3), 209-25.
- Grunert, Suzanne C., and Hans J. Juhl (1995), "Values, Environmental Attitudes, and Buying of Organic Foods," *Journal of Economic Psychology*, 16 (1), 39-62.
- Guido, Gianluigi (1991), "Implementing a Pan-European Marketing Strategy," *Long Range Planning*, 24 (5), 23-33.
- Gutman, Jonathan (1982), "A Means-End Model Based on Consumer Categorization Processes," *Journal of Marketing*, 46 (Spring), 60-72.
- Gutman, Jonathan (1991), "Exploring the Nature of Linkages Between Consequences and Values," *Journal of Business Research*, 22 (March), 143-8.
- Hassan, Salah S. and Erdener Kaynak (1994), "The Globalizing Consumer Market: Issues and Concepts," in *Globalization of Consumer Markets: Structures and Strategies*, Erdener Kaynak and Salah S. Hassan, eds. New York: International Business Press, 19-25.
- Hassan, Salah S. and Lea Prevel Katsanis (1994), "Global Market Segmentation Strategies and Trends," in *Globalization of Consumer Markets: Structures and Strategies*, Erdener Kaynak and Salah S. Hassan, eds. New York: International Business Press, 47-63.

- Helsen, Kristiaan, Kamel Jedidi, and Wayne S. DeSarbo (1993), "A New Approach to Country Segmentation Utilizing Multinational Diffusion Patterns," *Journal of Marketing*, 57 (October), 60-71.
- Henshel, Anne-Marie (1971), "The Relationship Between Values and Behavior: A Development Hypothesis," *Child Development*, 42, 1997-2007.
- Hoffman, Donna and Thomas P. Novak (1996), "Marketing in Hypermedia Computer-mediated Environments: Conceptual Foundations," *Journal of Marketing*, 60 (July), 50-68.
- Hofstede, Geert (1980), *Culture's Consequences: International Differences in Work Related Values*. Beverly Hills: Sage.
- Hoijtink, Herbert and Ivo W. Molenaar (1997), "A Multidimensional Item Response Model: Constrained Latent Class Analysis Using the Gibbs Sampler and Posterior Predictive Checks," *Psychometrika*, 62 (June), 171-89.
- Homer, Pamela M. and Lynn R. Kahle (1988). "A Structural Equation Analysis of the Value-Attitude-Behavioral Hierarchy," *Journal of Personality and Social Psychology*, 54 (April), 638-46.
- Jain, Subash C. (1989), "Standardisation of International Marketing Strategy: Some Research Hypotheses," *Journal of Marketing*, 53 (January), 70-9.
- Jain, Subash C. (1990), *International Marketing Management*. Boston: PWS.
- James, Don L., Richard M. Durand, and Robert A. Dreves (1976), "The Use of a Multi-Attribute Attitude Model in a Store Image Study," *Journal of Retailing*, 52, 23-32.
- Jedidi, Kamel, H.S. Jagpal, and Wayne S. DeSarbo, (1997), "STEMM: A General Finite Mixture Structural Equation Model," *Journal of Classification*, 14 (1), 23-50.
- Jolly, James P., Thomas J. Reynolds, and J.W. Slocum (1988), "Application of the Means-End Theoretic for Understanding the Cognitive Bases of Performance Appraisal," *Organizational Behavior and Human Decision Processes*, 41, 153-179.
- Kahle, Lynn R. (1983), *Social Values and Social Change: Adaptation to Life in America*. New York: Praeger.
- Kahle, Lynn R. (1986), "The Nine Nations of North America and the Value Basis of Geographic Segmentation," *Journal of Marketing*, 50 (April), 37-47.
- Kahle, Lynn R., Sharon E. Beatty, and Pamela Homer (1986), "Alternative Measurement Approaches to Consumer Values: The List of Values (LOV)

- and Values and Life Styles (VALS)," *Journal of Consumer Research*, 13, 405-9.
- Kale, Sudhir H. (1995), "Grouping Euroconsumers: A Culture-based Clustering Approach," *Journal of International Marketing*, 3 (3), 35-48.
- Kale, Sudhir H. and D. Sudharsan (1987), "A Strategic Approach to International Segmentation," *International Marketing Review*, 4 (Summer), 60-70.
- Kamakura, Wagner A. and José A. Mazzon (1991), "Value Segmentation: A Model for the Measurement of Values and Value Systems," *Journal of Consumer Research*, 18 (September), 208-18.
- Kamakura, Wagner A. and Thomas P. Novak (1992), "Value-System Segmentation: Exploring the Value of LOV," *Journal of Consumer Research*, 19 (June), 119-32.
- Kamakura, Wagner A., Thomas P. Novak, Jan-Benedict E.M. Steenkamp, and Theo M.M. Verhallen (1993), "Identification de Segments de Valeurs Pan-Européens par un Modèle Logit sur les Rangs avec Regroupements Successifs (Identifying Pan-European Value Segments With a Clusterwise Rank-Logit Model)," *Recherche et Applications en Marketing*, 8 (4), 30-55.
- Kamakura, Wagner A. and Gary J. Russel (1989), "A Probabilistic Choice Model for Market Segmentation and Elasticity Structure," *Journal of Marketing Research*, 26 (November), 379-90.
- Kass, Robert E. and Adrian E. Raftery (1995), "Bayes Factors," *Journal of the American Statistical Association*, 90 (June), 773-95.
- Kelly, George A. (1955), *The Psychology of Personal Constructs*. New York: W.W. Norton and Co. Inc.
- Kleinschmidt, Elko and Robert G. Cooper (1988), "The Performance Impact of International Orientation on Product Innovation," *European Journal of Marketing*, 22 (10), 56-71.
- Kotler, Philip (1997), *Marketing Management: Analysis, Planning, Implementation, and Control*. Upper Saddle River: Prentice-Hall, 9th ed.
- Kumar, Nirmalya (1997), "The Revolution in Retailing: from Market Driven to Market Driving," *Long Range Planning*, 30 (6), 830-5.
- Larntz, Kinley (1978), "Small-sample Comparisons of Exact Levels for Chi-squared Goodness-of-fit Statistics," *Journal of the American Statistical Association*, 73 (362), 253-63.
- Lenk, Peter J. and Wayne S. DeSarbo (1999), "Bayesian inference for finite mixtures of generalized linear models with random effects," *Psychometrika*, forthcoming.

- Lenk, Peter J., Wayne S. DeSarbo, Paul E. Green, and Martin R. Young (1996), "Hierarchical Bayes Conjoint Analysis: Recovery of Partworth Heterogeneity from Reduced Experimental Designs," *Marketing Science*, 15 (2), 173-91.
- Levitt, Theodore (1983), "The Globalisation of Markets," *Harvard Business Review*, May-June, 92-102.
- Lichtenstein, Donald R., Richard G. Netemeyer, and Scot Burton (1995), "Assessing the Domain Specificity of Deal Proness: A Field Study," *Journal of Consumer Research*, 22 (December), 314-26.
- Lindsay, Bruce, Clifford C. Clogg, and John Grego (1991), "Semiparametric Estimation in the Rasch Model and Related Exponential Response Models, Including a Simple Latent Class Model for Item Analysis," *Journal of the American Statistical Association*, 86 (413), 96-107.
- Louviere, Jordan J. and Richard D. Johnson (1991), "Using Conjoint Analysis to Measure Retail Image," *Research in Marketing*, Supplement 5, 137-56.
- M+M Eurodata (1996), *Structures, Chiffres d'Affaires et Enseignes de la Distribution Alimentaire et Non-Alimentaire Européenne*. Frankfurt: M+M.
- Mahajan, Vijay and Eitan Muller (1994), "Innovation Diffusion in a Borderless Global Market: Will the 1992 Unification of the European Community Accelerate Diffusion of New Ideas, Products, and Technologies?" *Technological Forecasting and Social Change*, 45, 221-35.
- Manzer, L. and S.J. Miller (1978), "An Examination of the Value-Attitude Structure in the Study of Donor Behavior," in *Proceedings of the American Institute of Decision Sciences*, Vol. 12, St. Louis, 532-38.
- Margolin, Barry H. and Richard J. Light (1974), "An Analysis of Variance for Categorical Data, II: Small Sample Comparisons with Chi Square and other Competitors," *Journal of the American Statistical Association*, 69, 755-64.
- Mazursky, David and Jacob Jacoby (1986) "Exploring the Development of Store Images," *Journal of Retailing*, 62 (2), 145-65.
- McCullagh, Peter, and John Ashford Nelder (1989), *Generalized Linear Models*, New York: Chapman and Hall.
- McLachlan, Geoffrey J., and Kaye E. Basford (1988), *Mixture Models*, New York: Marcel Dekker.
- Mittal, Banwari and Mysung-Soo Lee (1989), "A Causal Model of Consumer Involvement," *Journal of Economic Psychology*, 10, 363-89.

- Moskovitz, Howard R. and Samuel Rabino (1994), "Sensory Segmentation: An Organizing Principle for International Product Concept Generation," *Journal of Global Marketing*, 8 (1), 73-93.
- Nevin, John R. and Michael J. Houston (1980), "Image as a Component of Attraction to Intraurban Shopping Areas," *Journal of Retailing*, 56 (1), 77-93.
- Newell, Allen and Herbert A. Simon (1972), *Human Problem Solving*. Englewood Cliffs: Prentice-Hall.
- Neyman, Jerzy (1934), "On the Two Different Aspects of the Representative Method: The Method of Stratified Sampling and the Method of Purposive Selection," *Journal of the Royal Statistical Society*, 97, 558-606.
- Novak, Thomas P. and Bruce MacEvoy (1990), "On Comparing Alternative Segmentation Schemes: The List of Values (LOV) and Values and Life Styles (VALS)," *Journal of Consumer Research*, 17 (June), 105-09.
- Peter, J. Paul and Jerry C. Olson (1987), *Consumer Behaviour: Marketing Strategy Perspectives*. Homewood: Richard Irwin.
- Peter, J. Paul and Jerry C. Olson (1993), *Consumer Behavior and Marketing Strategy*. Homewood: Irwin, 3rd ed.
- Pieters, Rik, Hans Baumgartner and Doug Allen (1995), "A Means-end Chain Approach to Consumer Goal Structures," *International Journal of Research in Marketing*, 12 (3), 227-44.
- Pitts, Robert E. and Arch G. Woodside (1984), *Personal Values and Consumer Psychology*. Lexington: Lexington Books.
- Porter, Michael E. (1998), "Clusters and the New Economics of Competition," *Harvard Business Review*, November-December, 77-90.
- Poulsen, Carsten and Pierre Valette-Florence (1996), "Improvements in Means-end Research: A Heterogeneous Latent Markov Approach," *Paper presented at the Marketing Science Conference*, Gainesville, FL.
- Ramaswamy, Venkatram, Wayne S. DeSarbo, David J. Reibstein, and William T. Robinson (1993), "An Empirical Pooling Approach for Estimating Marketing Mix Elasticities with PIMS Data," *Marketing Science*, 12 (Winter), 103-24.
- Rand, W.M. (1971), "Objective Criteria for the Evaluation of Clustering Methods," *Journal of the American Statistical Association*, 85, 617-24.
- Rao, Vithala R. and Joel H. Steckel (1995), "A Cross-Cultural Analysis of Price Responses to Environmental Changes," *Marketing Letters*, 6 (1), 5-14.

- Rasch, G. (1960), *Probabilistic Models for Some Intelligence and Attainment Tests*. Copenhagen: Danmarks Paedagogiske Institut.
- Reynolds, Thomas J. and Alice B. Craddock (1988), "The Application of the MECCAS Model to the Development and Assessment of Advertising Strategy," *Journal of Advertising Research*, 28 (April/May), 43-54.
- Reynolds, Thomas J., Charles E. Gengler, and Daniel J. Howard (1995), "A Means-End Analysis of Brand Persuasion Through Advertising," *International Journal of Research in Marketing*, 12 (October), 257-66.
- Reynolds, Thomas J. and Jonathan Gutman (1984), "Advertising is Image Management," *Journal of Advertising Research* (February/March), 24, 27-37.
- Reynolds, Thomas J. and Jonathan Gutman (1988), "Laddering Theory, Method, Analysis, and Interpretation," *Journal of Advertising Research*, 28, 11-31.
- Reynolds, Thomas J. and Jerry C. Olson (1998), *The Means-End Approach to Understanding Consumer Decision Making: Applications to Marketing and Advertising Strategy*, Hillsdale: Lawrence Erlbaum.
- Richardson, Sylvia and Peter Green (1997), "On Bayesian Analysis of Mixtures with an Unknown Number of Components," *Journal of the Royal Statistical Society*, 59 (4), 731-92.
- Robert, Christian P. (1996), "Mixtures of Distributions: Inference and Estimation," in *Markov Chain Monte Carlo in Practice*, W.R. Gilks, S. Richardson, and D.J., Spiegelhalter, eds. New York: Chapman and Hall, 441-64.
- Rokeach, Milton J. (1973), *The Nature of Human Values*. New York: The Free Press.
- Rokeach, Milton J. and Sandra J. Ball-Rokeach (1989), "Stability and Change in American Value Priorities, 1968-1981," *American Psychologist*, 44 (5), 775-84.
- Ronen, Simcha and A.I. Kraut (1977), "Similarities among Countries based on Employee Work Values and Attitudes," *Columbia Journal of World Business*, 12 (2), 89-96.
- Ronen, Simcha and Oded Shenkar (1985), "Clustering Countries on Attitudinal Dimensions: A Review and Synthesis," *Academy of Management Review*, 10 (3), 435-54.
- Rossi, Peter E. and Greg M. Allenby (1993), "A Bayesian Approach to Estimating Household Parameters," *Journal of Marketing Research*, 30 (May), 171-82.

- Royall, Richard M. (1986), "Model Robust Confidence Intervals Using Maximum Likelihood Estimators," *International Statistical Review*, 54, 221-26.
- Samli, A. Coskun (1989), "Store Image Definition, Dimensions, Measurement, and Management," in *Retail Marketing Strategy*, A. Coskun Samli, ed. New York: Quorum.
- Scales, L.E. (1985), *Introduction to Non-linear Optimization*, London: MacMillan.
- Schank, Roger C. (1982), *Dynamic Memory: A Theory of Reminding and Learning in Computers and People*. New York: Cambridge University Press.
- Schwartz, G. (1978), "Estimating the Dimension of a Model," *Annals of Statistics*, 6, 461-64.
- Schwartz, Shalom H. (1992), "Universals in the Content and Structure of Values: Theoretical Advances and Empirical Tests in 20 Countries," in *Advances in Experimental Social Psychology*, Vol. 25, M. Zanna, ed. Orlando: Academic Press, 1-65.
- Schwartz, Shalom H. (1994), "Are There Universal Aspects in the Structure and Contents of Human Values?," *Journal of Social Issues*, 50 (4), 19-45.
- Sethi, S. Prakash (1971), "Comparative Cluster Analysis for World Markets," *Journal of Marketing Research*, 8 (August), 348-54.
- Shimp, Terence A. and Subash Sharma (1987), "Consumer Ethnocentrism: Construction and Validation of the CETSCALE," *Journal of Marketing Research*, 24 (August) 280-90.
- Shocker, Allan D. and V. Srinivasan (1979), "Multiattribute approaches for product concept evaluation and generation: A critical review," *Journal of Marketing Research*, 16 (May), 159-80.
- Singh, Surendra N. and Michael L. Rothschild (1983), "Recognition as a Measure of Learning from Television Commercials," *Journal of Marketing Research*, 20 (August), 235-248.
- Sirota, D. and J.M. Greenwood (1971), "Understand Your Overseas Work Force," *Harvard Business Review*, 49 (1), 53-60.
- Skinner, C.J. (1989), "Domain Means, Regression and Multivariate Analysis," in *Analysis of Complex Surveys*, Skinner, C.J., D. Holt, and T.M.F Smith, eds. New York: John Wiley & Sons, 59-87.
- Skinner, C.J., D. Holt, and T.M.F. Smith (1989), *Analysis of Complex Surveys*. New York: John Wiley & Sons.

- Sobel, Michael E. (1995), "The Analysis of Contingency Tables," in *Handbook of Statistical Modeling for the Social and Behavioral Sciences*, Gerhard Arminger, Clifford C. Clogg, and Michael E. Sobel, eds. New York: Plenum Press, 251-303.
- SRI Consulting (1999), *iVals: Internet VALS survey*, <http://future.sri.com/vals/iVALS.index.html>.
- Steenkamp, Jan-Benedict E.M. (1992), *De Europese Consument van Voedingsmiddelen in de Jaren '90*. Assen (NL): Van Gorcum.
- Steenkamp, Jan-Benedict E.M. and Hans Baumgartner (1998), "Assessing Measurement Invariance in Cross-National Consumer Research," *Journal of Consumer Research*, 25 (June), 78-90.
- Steenkamp, Jan-Benedict E.M., Frenkel ter Hofstede, and Michel Wedel (1999), "A Cross-National Investigation into the Individual and Cultural Antecedents of Consumer Innovativeness," *Journal of Marketing*, 63 (April), forthcoming.
- Steenkamp, Jan-Benedict E.M. and Hans C.M. Van Trijp (1996), "Task Experience and Validity in Perceptual Mapping: A Comparison of two Consumer-adaptive Techniques," *International Journal of Research in Marketing*, 13 (3), 265-76.
- Steenkamp, Jan-Benedict E.M. and Michel Wedel (1991), "Segmenting Retail Markets on Store Image Using a Consumer-Based Methodology," *Journal of Retailing*, 67 (3), 300-20.
- Steenkamp, Jan-Benedict E.M., Michel Wedel, George J. Avlonitis, Klaus G. Grunert, W. Bruce Traill, Frenkel ter Hofstede (1998), *A Consumer-led Approach to Foods in the EU*, Brussels: Commission of the European Communities.
- Stewart, David W. (1981), "The Application and Misapplication of Factor Analysis in Marketing Research," *Journal of Marketing Research*, 17 (February), 51-62.
- Sternquist, Brenda and Mudhao Kacker (1994), *European Retailing's Vanishing Borders*. London: Quorum.
- Szymanski, David M., Sundar G. Bharadwaj, and P. Rajan Varadarajan (1993), "Standardization versus Adaptation of International Marketing Strategy: An Empirical Investigation," *Journal of Marketing*, 57 (October), 1-17.
- Tackeuchi, Hirotaka and Michael E. Porter (1986), "Three Roles of International Marketing in Global Strategy," in *Competition in Global Industries*, Michael E. Porter ed. Boston: Harvard Business School Press, 111-46.

- Tanner, Martin and Wing Hung Wong (1987), "The Calculation of Posterior Distributions by Data Augmentation," *Journal of the American Statistical Association*, 82 (June), 528-40.
- Tellis, Gerard J. (1988), "The Price Elasticity of Selective Demand: A Meta-Analysis of Econometric Models of Sales," *Journal of Marketing Research*, 25 (November), 331-41.
- ter Hofstede, Frenkel, Anke Audenaert, Jan-Benedict E.M. Steenkamp, and Michel Wedel (1998), "An Investigation into the Association Pattern Technique as a Quantitative Approach to Measuring Means-End Chains," *International Journal of Research in Marketing*, 15 (January), 37-50.
- ter Hofstede, Frenkel, Jan-Benedict E.M. Steenkamp, and Michel Wedel (1999), "International Market Segmentation Based on Consumer-Product Relations," *Journal of Marketing Research*, 36 (February), 1-17.
- Thompson, M.E. (1997), *Theory of Sample Surveys*. New York: Chapman and Hall.
- Thurstone, L.L. (1927), "A Law of Comparative Judgment," *Psychological Review*, 34, 273-86.
- Titterton, D.M., A.F.M. Smith, and U.E. Makov (1985), *Statistical Analysis of Finite Mixture Distributions*. New York: John Wiley & Sons.
- Toler, C. (1975), "The Personal Values of Alcoholics and Addicts," *Journal of Clinical Psychology*, 31, 554-57.
- Tolman, Edward C. (1932), *Purposive Behaviour in Animals and Men*. New York: Century.
- Urban, Glen L. and John R. Hauser (1993), *Design and Marketing of New Products*, Englewood Cliffs: Prentice-Hall, 2nd ed.
- Usunier, Jean-Claude (1996), *Marketing across Cultures*, New York: Prentice Hall.
- Valette-Florence, Pierre and Bernard Rapacchi (1991), "Improvements in Means-end Chain Analysis using Graph Theory and Correspondence Analysis," *Journal of Advertising Research*, 31 (February/March), 30-45.
- Valli, Carlotta, Bruce Traill, and Rupert Loader (1997), "Analysis of the Yogurt Market in the European Union," *Report written for the Commission of the European Communities*, Brussels, Belgium.
- Vanden Abeele, Piet (1992), "A Means-end Study of Dairy Consumption Motivation," *Report written for the Commission of the European Communities*, EC regulation 1000/90 - 43 ST.

- Vandermerwe, Sandra and Marc-André L'Huillier (1989), "Euro-Consumers in 1992," *Business Horizons* (January/February), 34-40.
- Vriens, Marco, Michel Wedel, and Tom Wilms (1996), "Metric Conjoint Segmentation Methods: A Monte Carlo Comparison," *Journal of Marketing Research*, 33 (February), 73-85.
- Walker, Beth A. and Jerry C. Olson (1991), "Means-end Chains: Connecting Products with Self," *Journal of Business Research*, 22 (March), 111-18.
- Wedel, Michel and Wayne S. DeSarbo (1995), "A Mixture Maximum Likelihood Approach for Generalized Linear Models," *Journal of Classification*, 12 (1), 1-35.
- Wedel, Michel and Wayne S. DeSarbo (1996), "An Exponential Family Mixture Unfolding Methodology," *Journal of Business and Economic Statistics*, 14, 447-59.
- Wedel, Michel and Wagner A. Kamakura (1998), *Market Segmentation: Conceptual and Methodological Foundations*. Norwell: Kluwer Academic Publishers.
- Wedel, Michel, Wagner A. Kamakura, Wayne S. DeSarbo, and Frenkel ter Hofstede (1995), "Implications for Asymmetry, Nonproportionality, and Heterogeneity in Brand Switching from Piece-wise Exponential Mixture Hazard Models," *Journal of Marketing Research*, 32 (November), 457-62.
- Wedel, Michel, Frenkel ter Hofstede, and Jan-Benedict E.M. Steenkamp (1998), "Mixture Model Analysis of Complex Samples," *Journal of Classification*, 15 (2), 225-44.
- White, Halbert (1982), "Maximum Likelihood Estimation of Misspecified Models," *Econometrica*, 50 (January), 1-25.
- Wilkie, W.L. (1990), *Consumer Behavior*. New York: John Wiley & Sons, 2nd ed.
- Williams, D.A. (1976), "Improved Likelihood Ratio Tests for Complete Contingency Tables," *Biometrika*, 63 (1), 33-7.
- Wind, Yoram and Susan P. Douglas (1972), "International Market Segmentation," *European Journal of Marketing*, 6 (1), 17-25.
- Wind, Yoram and Vijay Mahajan (1997), "Issues and Opportunities in New Product Development: An Introduction to the Special Issue," *Journal of Marketing Research*, 34 (February), 1-12.
- Winston, Wayne L. (1987), *Operations Research: Applications and Algorithms*. Boston: PWS Publishers.

- Yavas, Ugur, Bronislaw J. Verhage, and Robert T. Green (1992), "Global Consumer Segmentation versus Local Market Orientation: Empirical Findings," *Management International Review*, 32 (3), 265-73.
- Yip, George S. (1995), *Total Global Strategy*. Englewood Cliffs: Prentice Hall.
- Young, Shirley and Barbara Feigin (1975), "Using the Benefit Chain for Improved Strategy Formulation," *Journal of Marketing*, 39 (July), 72-4.
- Zahn, D.A. and S.B. Fein (1979), "Large Contingency Tables with Large Cell Frequencies: A Model Search Algorithm and Alternative Measures of Fit," *Psychological Bulletin*, 86, 1189-200.

Index

A

Abbott, L., 43, 151
advertising, 1, 5, 10, 12, 44, 74, 92, 93,
95, 101, 110, 135, 136, 144, 145,
150
Agresti, Alan, 56, 151
Aitkin, Murray, 30, 79, 151
Akaike, H., 29, 79, 151, 152
Alba, Joseph W., 53, 151
Alden, Dana L., 2, 70, 151
Allen, Doug, 43, 159
Allenby, Greg M., 7, 77, 114, 151, 160
Amemiya, Takeshi, 33, 151
Aptech, 34, 151
Arora, Neeraj, 7, 77, 114, 151
Askegaard, Soren, 9, 15, 110, 151
attitudes, 8, 31, 40, 70, 95, 96
attributes, 3, 12, 16, 17, 41, 44, 45, 46,
47, 50, 51, 52, 53, 54, 55, 56, 57, 58,
60, 61, 62, 63, 65, 66, 67, 71, 73, 74,
85, 89, 92, 93, 99, 100, 101, 103,
111, 114, 115, 116, 125, 126, 130,
134, 139, 141, 144, 145
Audenaert, Anke, vii, 43, 162
Aurifeille, Jaques-Marie, 43, 152

Avlonitis, George J., 162

B

back translation, 14, 32, 86, 126, 141
Baker, Frank B., 75, 76, 152
Ball-Rokeach, Sandra J., 31, 160
Basford, Kaye E., 28, 38, 158
Bass, Frank M., 9, 152
Batagelj, Vladimir, 113, 154
Batra, Rajeev, 2, 70, 151
Baumgartner, Hans, 14, 43, 71, 75, 100,
142, 152, 159, 161
Bayes factor, 120, 127
Bayesian methods, 7, 18, 29, 38, 79,
107, 109, 111, 114, 117, 119, 120,
123, 127, 146, 148, 152, 153, 154,
155, 157, 158, 160, 179
Beatty, Sharon E., 31, 157
beef, 49, 50, 54, 57, 60, 62
Bell, Richard, 111, 124, 152
benefits, 3, 16, 17, 44, 45, 47, 50, 51,
52, 53, 55, 56, 58, 60, 61, 62, 63, 65,
66, 67, 70, 71, 73, 74, 85, 89, 92, 93,
99, 101, 103, 114, 144, 149
Bentler, Peter M., 57, 152
Best, N.G., 118, 152

Bettman, James, 44, 152
 Bharadwaj, Sundar G., 4, 162
 Binder, David A., 23, 152
 Blattberg, Robert C., 127, 152
 Bock, R. Darrell, 107, 152
 Bonett, Douglas G., 57, 152
 Boote, Alfred S., 10, 16, 73, 98, 152
 Bozdogan, Hamparsum, 29, 152, 153
 Brayne, C.E.G., 152
 Brislin, Richard W., 14, 32, 126, 153
 Burton, Scot, 95, 158

C

Carlin, Bradley P., 127, 153
 Carlin, John B., 155
 Carmone, Frank J., 113, 155
 Casella, George, 119, 153
 Cattin, Philippe, 153
 cell sparseness, 63
 Chen, Chuansheng, 71, 75, 153
 Claeys, C., 44, 153
 Clare, Donald A., 31, 153
 classification, 13, 38, 81, 83, 97, 98, 112, 126
 Clogg, Clifford C., 76, 158, 161
 cluster analysis, 6, 8, 9, 11, 12, 18, 97, 98, 99, 100, 112, 113, 114, 146
 cluster sample, 27
 Cochran, W.G., 22, 27, 153
 communication, 7, 16, 17, 71, 73, 85, 90, 99, 100, 101, 144
 communication strategies, 16, 17, 71, 73, 144
 complex samples, 25, 29
 conditional independence, 50, 51, 53, 55, 56, 57, 59, 60, 61, 62, 65, 140
 content analysis, 48, 60
 convergent validity, 17, 45, 50, 51, 52, 54, 55, 57, 59, 62, 63, 64, 65, 147
 Cooper, Robert G., 70, 153, 157
 correspondence analysis, 97
 cost effectiveness, 16, 111
 Craddock, Alice B., 44, 159
 Craig, C. Samuel, 154
 Cron, W.L., 114, 153
 culture, 3, 8, 9, 14, 31, 72, 102, 110, 129, 132, 141

D

Danes, Jeffrey, 153
 data augmentation, 115
 data collection, 15, 16, 46, 47, 48, 51, 53, 67, 71, 74, 85, 146
 Davies, Ross, 111, 124, 152
 Day, Ellen, 8, 9, 153
 Dempster, A.P., 28, 153
 DeSarbo, Wayne S., 6, 7, 9, 13, 16, 24, 56, 72, 110, 113, 114, 140, 153, 156, 157, 158, 159, 163
 Diebolt, J., 7, 154
 Dillon, William R., 6, 66, 154
 discriminant analysis, 97
 distribution, 1, 5, 110, 111, 130, 133, 140, 145
 diversification, 110
 Douglas, Susan P., 9, 11, 14, 32, 109, 110, 140, 152, 154, 164
 Dreves, Robert A., 125, 156
 Durand, Richard M., 125, 156

E

Edwards, Elizabeth A., 114, 153
 entropy, 79, 80, 87
 entry strategy, 6, 111
 Eshgi, Abdolreza, 9, 154
 European Union, 5, 10, 18, 69, 72, 85, 102, 110, 111, 112, 125, 135, 139, 145, 149, 162, 163
 Eurostat, 126, 135, 154
 exponential family, 6, 21, 24, 26, 41, 140

F

Fahrmeier, Ludwig, 26, 154
 Farley, John U., 4, 154
 Feather, N.T., 31, 154
 Feigin, Barbara, 46, 74, 164
 Fein, S.B., 65, 164
 Ferligoj, Anuska, 113, 154
 Fienberg, S., 60, 154
 Fox, Richard J., 8, 9, 153
 Frank, Ronald E., 11, 152, 154, 155
 Frederick, Donald G., 66, 154
 full-conditional distribution, 7, 119, 120, 140

G

Gelfand, Alan E., 119, 153, 155
 Gelman, Andrew, 7, 120, 148, 155
 Geman, Donald, 117, 155
 Geman, Stuart, 117, 155
 generalized linear model, 6, 24, 140, 157
 Gengler, Charles E., 101, 160
 geodemographics, 135, 136, 139
 George, Edward I., 119, 127, 152, 153, 157, 162, 164
 Ghose, Sanjay, 97, 155
 Ghosh, Malay, 76, 155
 Gibbs-sampler, 117, 118, 119, 120, 135, 137, 140, 153, 155, 156
 Ginter, James L., 7, 77, 114, 151
 globalization, 2, 4, 69, 109
 Goldstein, Matthew, 6, 154
 Goodman, L.A., 62, 155
 goodness-of-fit, 57
 Gordon, A.D., 112, 155
 Green, Paul E., 66, 113, 155, 158
 Green, Peter, 160
 Green, Robert T., 11, 15, 16, 73, 98, 164
 Greenwood, J.M., 8, 151, 161
 Grego, John, 76, 158

Grenander, U., 119, 155
 Griffin, Abbie S., 101, 155
 Grunert, Klaus G., 44, 46, 47, 52, 155, 162
 Grunert, Suzanne C., 31, 44, 46, 52, 155
 Guido, Gianluigi, 40, 155
 Gutman, Jonathan, 44, 45, 46, 47, 51, 55, 71, 73, 74, 75, 101, 155, 156, 160

H

Hassan, Salah S., 2, 10, 70, 156
 Hauser, John R., 74, 101, 163
 Helsen, Kristiaan, 9, 13, 16, 72, 110, 156
 Henshel, Anne-Marie, 31, 156
 heterogeneity, 1, 3, 11, 14, 67, 72, 73, 76, 77, 79, 89, 98, 100, 111, 114, 121, 137, 138, 139, 145, 179
 heuristic, 6, 7, 8, 11, 12, 13, 16, 18, 30, 146
 hierarchical structure, 44, 115, 118, 148
 Hildebrand, David K., 56, 153
 Hoffman, Donna, 150, 156
 Hofstede, Geert, 8, 156
 Hoijsink, Herbert, 7, 156
 holdout samples, 97, 100
 Holt, D., 22, 78, 161
 Homer, Pamela M., 31, 93, 156, 157
 Houston, Michael J., 124, 159
 Howard, Daniel J., 101, 160
 Howard, Elizabeth, 111, 124, 152
 Huszagh, Sandra M., 8, 9, 153
 Hutchinson, J. Wesley, 53, 151

I

information criteria, 87, 89, 120, 147
 international strategy, 2, 8
 Internet, 150

item response theory, 75, 76, 148

J

Jacoby, Jacob, 125, 158
 Jagpal, H.S., 25, 156
 Jain, Subash C., 14, 32, 40, 73, 109, 156
 James, Don L., 125, 151, 152, 156
 Jedidi, Kamel, 9, 13, 16, 25, 72, 110, 156
 Johnson, Richard D., 125, 158
 Jolly, James P., 31, 156
 Juhl, Hans J., 31, 155

K

Kacker, Mudhao, 124, 162
 Kahle, Lynn R., 8, 31, 32, 66, 86, 93, 110, 156, 157
 Kale, Sudhir H., 11, 72, 98, 157
 Kamakura, Wagner A., 3, 4, 10, 13, 15, 16, 23, 31, 33, 39, 40, 73, 74, 78, 86, 95, 98, 140, 157, 163
 Kass, Robert E., 120, 157
 Katsanis, Lea Prevel, 2, 10, 70, 156
 Kaynak, Erdener, 70, 154, 156
 Kelly, George A., 54, 66, 157
 Kleinschmidt, Elko, 70, 157
 K-means, 11, 18, 98, 99, 100
 Kotler, Philip, 100, 157
 Kraut, A.I., 8, 160
 Kumar, Nirmalya, 124, 157

L

laddering, 17, 45, 46, 47, 48, 50, 51, 52, 53, 54, 55, 56, 57, 59, 60, 61, 62, 63, 64, 65, 66, 85, 147
 Laird, N.M., 28, 153
 Larntz, Kinley, 59, 61, 63, 157
 Lee, Mysung-Soo, 95, 158

Lee, Shin-ying, 71, 75, 153
 Lehmann, Donald R., 154
 Lenk, Peter J., 7, 157, 158
 Levitt, Theodore, 2, 4, 158
 Lichtenstein, Donald R., 95, 158
 Lieberman, Murray, 107, 152
 Light, Richard J., 59, 158
 Lindsay, Bruce, 76, 158
 List of Values, 32, 33, 34, 66, 86, 157, 159
 Loader, Rupert, 85, 163
 Louis, Thomas A., 127, 153, 158
 Louviere, Jordan J., 125, 158

M

M+M Eurodata, 125, 158
 MacEvoy, Bruce, 31, 159
 macro-level, 7, 11
 Madsen, Tage Koed, 9, 15, 110, 151
 Mahajan, Vijay, 2, 70, 113, 153, 158, 164
 Makov, U.E., 24, 28, 30, 75, 162
 Manzer, L., 31, 158
 Margolin, Barry H., 59, 158
 Massy, William F., 11, 154
 maximum likelihood, 23, 25, 26, 27, 28, 29, 30, 32, 34, 35, 36, 37, 38, 40, 78, 100, 146, 148
 Mazursky, David, 125, 158
 Mazzon, José A., 10, 23, 31, 33, 95, 157
 McCullagh, Peter, 24, 158
 McLachlan, Geoffrey J., 28, 38, 158
 MCMC methods, 7, 149
 means-end chains
 association pattern technique, 147
 laddering, 17, 45, 46, 47, 48, 50, 51, 52, 53, 54, 55, 56, 57, 59, 60, 61, 62, 63, 64, 65, 66, 85, 147
 MEC map, 90, 97

theory, 6, 16, 17, 22, 32, 33, 43, 44,
45, 51, 65, 71, 73, 102, 116, 144,
145, 147
means-end chains, 16, 17, 44, 45, 46, *
47, 48, 50, 51, 52, 53, 57, 58, 59, 63,
65, 66, 67, 71, 73, 74, 80, 89, 90, 92,
93, 97, 99, 101, 103, 144, 147, 149
measurement instrument, 14, 15, 17,
18, 135, 140, 141
media consumption, 18, 95, 135, 139,
144
Miller, S.J., 31, 158
Mittal, Banwari, 95, 158
mixture models, 6, 7, 10, 13, 17, 21, 22,
23, 24, 25, 26, 28, 29, 33, 89, 113,
114, 146, 179
model fit, 57, 64, 98
model-based, 6, 9, 13, 16, 25, 144, 145,
150
Molenaar, Ivo W., 7, 156
Monte Carlo study, 18, 80
Moskovitz, Howard R., 73, 158
Muller, Eitan, 2, 70, 158
multi-attribute model, 114, 115, 125

N

Nelder, John Ashford, 24, 140, 158
Netemeyer, Richard G., 95, 158
Nevin, John R., 124, 159
Newell, Allen, 71, 73, 159
Neyman, Jerzy, 22, 159
North Atlantic Free Trade Area, 5, 69
Novak, Thomas P., 31, 39, 40, 74, 86,
150, 156, 157, 159
NUTS2, 126

O

olive oil, 54, 60
Olson, Jerry C., 43, 46, 51, 65, 67, 71,
73, 159, 160, 163

P

pan-European, 10, 32, 40, 85, 86, 88,
93, 98, 100, 125, 126, 144, 179
pan-regional, 2, 32, 70, 71, 72, 73, 80,
101
personality, 18, 31, 85, 95, 144
Peter, J. Paul, 43, 51, 67, 73, 151, 152,
157, 158, 159, 160
pick-any, 17, 97
Pieters, Rik, 43, 159
Pitts, Robert E., 73, 159
Porter, Michael E., 5, 110, 159, 162
positioning, 1, 90, 92, 99, 101, 124, 130
posterior distribution, 7, 118, 119, 120,
127, 130
posterior probability, 38, 79, 107, 129
Poulsen, Carsten, 67, 159
predictive validity, 18, 72, 94, 97
price, 112, 128, 130
probability distribution, 6, 7, 25, 26, 33,
57, 64, 75, 77, 78, 80, 81, 83, 84,
117, 119, 122, 140
product development, 12, 16, 17, 44,
71, 73, 74, 85, 92, 93, 99, 100, 101,
144
profiles, 40, 47, 93, 95, 96, 135, 137,
150
promotions, 5, 10, 95, 110
pseudo information criteria, 79, 87, 88,
89, 147
pseudo likelihood, 30, 98
pseudo maximum likelihood, 23, 24,
25, 26, 27, 28, 29, 30, 32, 34, 35, 36,
37, 38, 39, 41, 72, 78, 100, 108, 146,
148
purchasing power, 2, 70, 136

R

Rabino, Samuel, 12, 16, 73, 158
Raftery, Adrian E., 120, 157
Ramaswamy, Venkatram, 79, 159

- Rand, W.M., 98, 159
 Rao, Vithala R., 128, 159
 Rapacchi, Bernard, 43, 163
 Rasch, G., 76, 155, 158, 159
 regression, 6, 21, 24, 113, 115, 127
 Reibstein, David J., 159
 response tendencies, 14, 16, 18, 71, 75,
 76, 79, 89, 90, 98, 99, 100, 144, 148
 restrictions, 7, 18, 78, 100, 111, 112,
 113, 114, 116, 117, 120, 124, 127,
 137, 138, 139, 140, 141, 145, 146
 retailing, 6, 18, 111, 124, 125, 130, 135,
 137, 139, 145
 Reynolds, Thomas J., 31, 44, 45, 46,
 71, 74, 101, 156, 159, 160
 Richardson, Sylvia, 160
 Robert, Christian P., 7, 115, 154, 160
 Robinson, William T., 159
 Rokeach, Milton J., 10, 31, 74, 160
 Ronen, Simcha, 8, 15, 160
 Rossi, Peter E., 7, 151, 160
 Rothschild, Michael L., 53, 161
 Royall, Richard M., 28, 29, 160
 Rubin, Donald B., 28, 30, 79, 151, 153,
 155
 Russel, Gary J., 140, 157
- S**
- Samli, A. Coskun, 124, 160
 sampling design, 15, 16, 17, 22, 23, 24,
 25, 26, 27, 29, 32, 36, 40, 80, 98,
 100, 144, 146, 148
 Sanford, Donald G., 153
 Scales, L.E., 28, 153, 160
 Schank, Roger C., 44, 160
 Schwartz, G., 29, 160
 Schwartz, Shalom H., 10, 30, 31, 161
 segmentation
 constrained segmentation, 112
 cross-national segmentation, 32
 domestic segmentation, 72, 100, 148
 geographic segmentation, 13, 114,
 118, 119, 127, 135
 international segmentation, 1, 3, 4,
 6, 7, 9, 10, 13, 14, 15, 16, 17, 18,
 24, 41, 45, 67, 72, 73, 74, 79,
 102, 110, 143, 144, 146, 147,
 148, 149, 150
 response-based segmentation, 110,
 114, 139
 segmentation basis, 3, 7, 9, 14, 16,
 99, 111, 112, 143
 value segmentation, 10, 17, 30, 32,
 40, 147
 segmentation, 1, 2, 3, 4, 5, 6, 7, 8, 9, 10,
 11, 12, 13, 14, 15, 16, 17, 18, 21, 22,
 23, 26, 31, 40, 48, 72, 73, 74, 75, 76,
 80, 87, 98, 99, 100, 111, 112, 113,
 114, 119, 124, 125, 138, 139, 140,
 143, 144, 145, 146, 147, 148, 149,
 150, 179
 Sethi, S. Prakash, 7, 8, 10, 110, 161
 Sharma, Subash, 95, 161
 Shenkar, Oded, 8, 15, 160
 Sheth, Jagdish N., 9, 154
 Shimp, Terence A., 95, 161
 Shocker, Allan D., 73, 161
 Simon, Herbert A., 71, 73, 159
 Singh, Surendra N., 53, 161
 Sirota, D., 8, 161
 Skinner, C.J., 22, 23, 26, 30, 78, 161
 Slocum, J.W., 31, 156
 Smith, Adrian F.M., 24, 28, 30, 75,
 119, 155, 162
 Smith, T.M.F., 22, 78, 161
 Sobel, Michael E., 62, 161
 spatial segments, 111
 Spiegelhalter, D.J., 152, 160
 Srinivasan, V., 66, 73, 155, 161
 standardization, 2, 4, 11, 48, 69, 147
 Steckel, Joel H., 128, 159
 Steenkamp, Jan-Benedict E.M., vii, 2,
 9, 14, 21, 31, 43, 47, 69, 70, 71, 73,

75, 95, 100, 114, 125, 142, 149, 151,
152, 157, 161, 162, 163
Stern, Hal S., 155
Sternquist, Brenda, 124, 162
Stevenson, Harold W., 71, 75, 153
Stewart, David W., 6, 162
store image, 111, 114, 124, 125, 126,
130, 134, 135, 139, 145
stratification, 15, 22, 35, 148
Sudharsan, D., 11, 157
Swinnen, A., 44, 153
Szymanski, David M., 4, 162

T

Tackeuchi, Hirotaka, 5, 110, 162
Tangpanichdee, Vanchai, 66, 154
Tanner, Martin, 115, 162
target market selection, 6, 100, 109,
111, 112, 145, 149
target markets, 6, 109, 111
Tellis, Gerhard J., 128, 162
ter Hofstede, Frenkel, 14, 21, 31, 43,
69, 71, 73, 95, 161, 162, 163, 179
Thomas, A., 152, 153, 156, 157, 159,
160
Thompson, M.E., 22, 162
Thurstone, L.L., 33, 162
Titterington, D.M., 24, 28, 30, 75, 162
Toler, C., 31, 162
Tolman, Edward C., 43, 162
Traill, W. Bruce, 85, 162, 163
Tutz, Gerhard, 26, 154

U

Urban, Christine, 9, 154
Urban, Glen L., 74, 101, 163
Usunier, Jean-Claude, 8, 163

V

Valette-Florence, Pierre, 43, 67, 152,
159, 163
Valli, Carlotta., 85, 163
values, 3, 4, 10, 16, 17, 18, 23, 28, 30,
31, 32, 33, 34, 37, 39, 40, 41, 44, 45,
48, 50, 51, 52, 53, 55, 56, 57, 58, 60,
61, 62, 63, 65, 66, 67, 70, 71, 73, 74,
79, 80, 84, 86, 87, 88, 89, 92, 93, 99,
100, 101, 103, 116, 120, 122, 123,
128, 135, 138, 144
Van Trijp, Hans C.M., 47
Vanden Abeele, Piet, 44, 47, 54, 153,
163
Vandermerwe, Sandra, 9, 163
Varadarajan, P. Rajan, 4, 162
Verhage, Bronislaw J., 11, 15, 16, 73,
98, 164
Verhallen, Theo M.M., 157
Vriens, Marco, 114, 163

W

Wachspress, David P., 113, 155
Walker, Beth A., 46, 51, 65, 163
Wedel, Michel, vii, 3, 4, 6, 14, 21, 24,
31, 43, 69, 71, 73, 78, 95, 98, 113,
114, 125, 161, 162, 163
White, Halbert, 28, 29, 163
Wilkie, W.L., 110, 163
Williams, D.A., 60, 62, 164
Wilms, Tom, 114, 163
Wind, Yoram, 11, 70, 110, 140, 154,
164
Winston, Wayne L., 116, 164
within-segment heterogeneity, 7, 77,
116, 121, 123, 138, 139, 146
Wong, Wing Hung, 115, 162
Woodside, Arch G., 73, 159

Y

Yavas, Ugur., 11, 15, 16, 73, 98, 164

Yip, George S., 2, 5, 69, 110, 164

yogurt, 54, 60, 85, 86, 87, 90, 92, 93,
95, 96, 100, 104, 144

Young, Martin R., 158

Young, Shirley, 46, 74, 164

Z

Zahn, D.A., 65, 164

Samenvatting

Dit proefschrift behandelt een aantal methodologische en inhoudelijke aspecten van internationale marktsegmentatie. Hoofdstuk 1 geeft een overzicht van voorgaand onderzoek op het gebied van internationale marktsegmentatie. De segmentatie studies zijn geïdentificeerd aan de hand van een drietal dimensies die samenhangen met de effectiviteit van internationale marktsegmentatie. Op basis van een kritische evaluatie van deze studies is een aantal methodologische en inhoudelijke issues geïdentificeerd die verdere aandacht behoeven. Deze issues zijn in de hoofdstukken 2 tot en met 5 nader onderzocht en methodologieën zijn ontwikkeld die dienen als oplossing voor deze problemen.

Hoofdstuk 2 behandelt een belangrijk methodologisch probleem in internationale marktsegmentatie. Internationale steekproeven worden in de regel verzameld volgens specifieke steekproef designs. Vaak wordt de steekproefgrootte gelijk gehouden over landen of wordt een minimale steekproefgrootte per land gespecificeerd om additionele analyses per land toe te laten. De grootte van nationale steekproeven is veelal niet proportioneel aan de populatiegrootte hetgeen leidt tot een impliciete stratificatie van de steekproef per land. Bovendien wordt in sommige gevallen het steekproef design per land aangepast. Huidige methoden voor zowel nationale als internationale segmentatie houden geen rekening met dergelijke designs, hetgeen kan leiden tot resultaten die niet

representatief zijn voor de populatie onder studie en tot onjuiste beslissingen van managers.

In hoofdstuk 2 wordt een methode voorgesteld, pseudo maximum likelihood (PML), om mengsel segmentatie modellen te schatten indien steekproeven verzameld zijn volgens een complex steekproef design. De procedure identificeert consistente schattingen van de model parameters. Daarbij worden gemodificeerde of pseudo-informatie criteria voorgesteld om op correcte wijze het aantal segmenten te bepalen. De methode is toegepast in een internationale segmentatiestudie op basis van persoonlijke waarden. Pseudo maximum likelihood schattingen zijn vergeleken met traditionele maximum likelihood schattingen die geen rekening houden met het steekproef design. De resultaten tonen aan dat schattingen van segment groottes en segment-specifieke parameters een bias kunnen vertonen indien geen rekening wordt gehouden met het internationale design van steekproeftrekken. Verder laat de empirische toepassing zien dat het gebruik van traditionele informatie criteria kan leiden tot onjuiste conclusies omtrent het aantal segmenten. Dit betekent dat het toepassen van traditionele methoden in internationale marktsegmentatie kan leiden tot incorrecte conclusies en onjuiste management beslissingen.

Hoofdstuk 3 ontwikkelt en valideert een kwantitatief instrument voor het meten van means-end chains (middel-doel ketens). Dit hoofdstuk dient hoofdstuk 4, waar een methodologie wordt ontwikkeld voor het identificeren van internationale segmenten op basis van means-end chains. De traditionele methode voor het meten van means-end chains, laddering, is niet geschikt voor het verzamelen van grote en representatieve internationale steekproeven. In hoofdstuk 3 wordt een kwantitatief meetinstrument voorgesteld voor het meten van means-end chains. Dit instrument wordt aangeduid als de association pattern technique (APT; associatie patroon techniek). De APT methode biedt de mogelijkheid om means-end chain data te verzamelen middels grote en representatieve steekproeven en is een geschikte meettechniek voor toepassingen van internationale marktsegmentatie.

Ten aanzien van de validiteit van APT zijn twee centrale vraagstukken onderzocht. Ten eerste meet APT de verbanden tussen

means-end chain concepten (attributen, benefits, en waarden) in twee afzonderlijke taken. In statistische termen betekent dit dat attributen en waarden conditioneel onafhankelijk worden verondersteld, gegeven benefits. Ten tweede is de convergente validiteit onderzocht tussen APT en de traditionele laddering methode. In hoofdstuk 3 zijn deze vraagstukken getoetst in een empirische studie. Consistente ondersteuning is gevonden voor de conditionele onafhankelijkheid van attributen en waarden, gegeven benefits, voor vier productcategorieën. Statistische toetsen voor convergente validiteit van APT en laddering tonen aan dat de technieken gelijke means-end structuren identificeren. Dit betekent dat APT een valide techniek is voor het meten van means-end chains en geschikt is voor het identificeren van internationale marktsegmenten.

Voortbouwend op hoofdstukken 2 en 3 wordt in hoofdstuk 4 een geïntegreerde methodologie van dataverzameling en analysemethode voor internationale segmentatie voorgesteld, gebruik makend van means-end chain theorie. De methodologie maakt het mogelijk transnationale segmenten te identificeren die bestaan uit consumenten met gelijke means-end chain structuren. De means-end chain theorie bouwt voort op eerder werk van economen en psychologen die inzagen dat consumenten preferenties niet alleen bepaald worden door de concrete produktkarakteristieken, maar ook vanwege de betekenis van deze producten voor deze consumenten. Means-end chain theorie verbindt het produkt met de consument door een hiërarchische structuur op te leggen aan produkt attributen, de voordelen van produkt gebruik (benefits) en persoonlijke waarden van consumenten. De methodologie biedt de mogelijkheid strategieën voor produkt ontwikkeling te combineren met communicatiestrategieën op internationaal segment niveau en kan gebruikt worden om producten en communicatie in overeenstemming te brengen met de verlangens van internationale segmenten van consumenten.

Een internationaal segmentatiemodel is ontwikkeld dat is afgestemd op data die verzameld zijn middels de APT methode. Het model houdt rekening met een aantal zaken die van belang zijn voor het op een valide wijze identificeren van internationale segmenten. Het

model is een eindig mengsel model dat segmenten identificeert op basis van de gelijkheid van probabilistische relaties tussen produkt attributen, benefits, en persoonlijke waarden van consumenten. Aangezien het meetinstrument gevoelig kan zijn voor verschillen in respons gedrag en omdat deze verschillen de identificatie van transnationale segmenten in de weg kunnen staan, corrigeert het model voor dergelijke verschillen, zowel tussen als binnen landen. Een specificatie van het model is ontwikkeld waarin drempel parameters de responsgeneigdheid beschrijven en die overeenkomsten vertoont met het Rasch model. De segment proporties zijn gerelateerd aan landen door middel van een 'concomittant variables'-specificatie zodat segmentgroottes kunnen variëren over landen. Dit maakt het mogelijk globale of pan-regionale segmenten te identificeren.

De internationale segmentatie methodologie is toegepast op een uitgebreide data set verzameld middels een omvangrijke steekproef van huishoudens in elf landen van de Europese Unie. Vier segmenten werden geïdentificeerd, waarvan één pan-Europees segment en drie transnationale segmenten. De segmenten vertegenwoordigden verschillende means-end chain structuren en uit de patronen tussen attributen, benefits, en waarden volgden strategische implicaties voor produkt ontwikkeling en communicatie. De segmenten toonden significante relaties met socio-demografische kenmerken, consumptiepatronen, mediaconsumptie en persoonlijkheids-karakteristieken. Dit draagt bij tot de identificeerbaarheid en toegankelijkheid van de segmenten. Additionele analyses ondersteunden de validiteit van het model. De performance van het model is verder onderzocht door middel van een Monte Carlo studie en de resultaten tonen aan dat het model goed in staat is de werkelijke parameters te identificeren, gegeven de typische structuur van APT data. Een hoge predictieve validiteit van het model is aangetoond en een empirische vergelijking wees uit dat het model niet te benaderen valt met de traditionele K-means clustering techniek.

In hoofdstuk 5 wordt een tweede methode voorgesteld ter bevordering van de geografische configuratie van segmenten. De methode is vooral nuttig voor de selectie en afbakening van nieuwe

markten door internationaal expanderende bedrijven. Eerder voorgestelde methoden van internationale segmentatie leiden niet noodzakelijkerwijs tot segmenten die toegankelijk zijn met efficiënte logistieke activiteiten. Vooral indien de kosten van fysieke distributie een belangrijk onderdeel uitmaken van de totale kosten van productie en marketing, is het belangrijk dat een segment een aaneengesloten gebied beslaat en niet bestaat uit verspreide segmenten die kunnen worden geïdentificeerd met bestaande segmentatie methoden. Hoofdstuk 5 ontwikkelt een flexibel internationaal segmentatie model dat aaneengesloten geografische segmenten identificeert aan de hand van gegevens van consumenten. Een Bayesiaanse model-specificatie is gekozen die de structuur van het probleem goed beschrijft. Het model is gebaseerd op multi-attribuu theorie van preferentie formatie en corrigeert voor de heterogeniteit binnen segmenten. Een 'Gibbs sampling' procedure is ontwikkeld voor het schatten van het model.

Het model is geïllustreerd aan de hand van een internationale segmentatie studie op basis van winkel imago. Binnen de internationale retailing vormt geografische expansie een belangrijke groeistrategie. Op basis van belangen die consumenten toekennen aan verschillende attributen van winkel imago zijn vijf segmenten geïdentificeerd in zeven landen van de Europese Unie. De segmenten waren onderscheidend in de patronen van attribuut belangen van winkel imago. De resultaten illustreerden de mogelijkheden voor internationaal expanderende retailers om geografische gebieden te bedienen met passende winkel imago formules. De resultaten toonden ook de toegankelijkheid van de segmenten door middel van communicatie en logistieke operaties. Een vergelijking van het voorgestelde model met een ongerestricteerde variant van het model, waarbij geen restricties op de configuratie van de segmenten werd opgelegd, liet weinig verschillen zien in de heterogeniteit binnen segmenten. Dit betekent dat de voordelen van het opleggen van geografische restricties de 'responsiveness' van segmenten nauwelijks beïnvloedt. In de discussie van hoofdstuk 5 worden uitbreidingen van het model gegeven voor diverse segmentatiebases en voor een brede klasse van restricties die opgelegd kunnen worden aan de segmenten.

In hoofdstuk 6 volgt een integrale discussie van de verschillende studies beschreven in dit proefschrift. Ter afsluiting wordt een aantal beperkingen van dit proefschrift beschreven, alsmede suggesties voor verder onderzoek.

Curriculum Vitae

Frenkel ter Hofstede was born in Oldenzaal (the Netherlands) on September 18, 1968. He received his masters degree in 1994 at the University of Groningen, with a major in econometrics and a minor in marketing and his masters thesis dealt with brand choice and purchase timing decisions of consumers. After graduation he married Judith J.A. Stomp and started a joint appointment at Wageningen University and the University of Groningen as a Ph.D. student and research associate. He was a member of an international research consortium that studied the development of pan-European marketing strategies for food companies, financed by the Commission of the European Communities. As from August 1999 he will start as an Assistant Professor of Marketing at the Graduate School of Industrial Administration of Carnegie Mellon University.

His research interests lie in developing and applying advanced marketing models and techniques to substantive marketing problems, particularly in the area of international marketing and market segmentation. His methodological interest is in hazard models and methods for describing discrete and continuous consumer heterogeneity, including mixture models, random effects modeling, and hierarchical Bayes. He published in a number of academic journals, including *Economics letters*, *International Journal of Research in Marketing*, the *Journal of Classification*, the *Journal of Marketing Research*, and the *Journal of Marketing*.

Printed by Universal Press