

Marketing Engineering 2020-2021

Estelle Severs - r0793932

Intro: Wat is Marketing Engineering

Aan de hand van modellen de realiteit op een eenvoudige manier kunnen visualiseren of voorstellen. -> Dit is de kern van Marketing Engineering

Rest van de les is vooral introductie en niet echt van belang voor het examen.

Informatieverzameling

Waar halen we onze info vandaan?

- Primaire Bronnen
 - Resultaten van eigen onderzoek
- Secundaire Bronnen
 - Resultaat van eerder uitgevoerd onderzoek
 - _ Intern: Gegevens die binnen een bedrijf worden bijgehouden
 - _ Extern: Toegankelijk voor iedereen vs toegankelijk op voorwaarde van betaling.

1. Primaire Bronnen

Resultaat van eigen onderzoek (of met behulp van onderzoeksbureaus)

- *Exploreren*
 - Vrijblijvend inzichten leveren en ontwikkelen
 - Minder gestructureerd
 - Kwalitatief
- *Confirmeren*
 - Concrete vraag beantwoorden, een beslissing helpen nemen
 - Meer gestructureerd
 - Kwantitatief

Kwalitatief onderzoek

Deze onderzoeken gebeuren meestal onder de vorm van *Diepte interviews* (*Intensief individueel interview*) of *Groepsgesprekken* (*Focus group*).

Hier staan de WAAROM en HOE centraal. We zoeken *Relevante* personen (met een sterke mening of met ervaring) om hun mening te uiten over het product.

Inhoud: Analyseren en verwerken van woord- en zinnen materiaal (bv tellen hoe vaak een thema aan bod komt, individueel of in combo met een ander thema)

Resultaat: 'Het in kaart brengen (van een bepaald beleavingsgebied. Grote thema's, onderliggende assen, tegenstellingen en conflicten, etc..

Basisinzichten -> Toetsbare hypothesen.

Grappig voorbeeldje: <https://www.youtube.com/watch?v=tVq1wglN62E>

Kwantitatief onderzoek

Belangrijk hier is:

- Sterke structurering (precieze antwoorden vragen dat men kan coderen of kwantificeren)
- Voldoende omvang
- Representatief (vaak een groot kostelijk onderzoek)

We kunnen dit onderverdelen in 3 soorten:

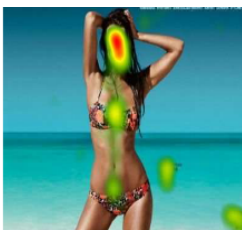
I. Observatie

= Bekijken en coderen van gedrag om systematische patronen te ontdekken.

Probleem: Niet vaak sociaal wenselijk, mensen gaan reageren op het feit dat ze geobserveerd worden. We zoeken dus varianten die niet opvallen.

-> Er bestaan niet-indringende varianten (daarom niet altijd legaal):

- *Garbology* (zoeken wat mensen weggooien in hun afval)
- *Scanner data* (zoals Nielsen kind of)
- *Eye-tracking* (waar kijken mensen naar bij reclame, is er niets te afleidend?)
- *Clickstream-data*



Voorbeeld Eye-tracking:

De mensen kijken duidelijk te veel naar het gezicht van deze persoon (en andere delen van het lichaam) terwijl ze moeten kijken naar het logo van de reclame (rechtsonder). Dit zou slechte reclame zijn.

II. Experiment

= Observeren wat mensen doen maar onder gecontroleerde omstandigheden.

- ★ *Manipulatie van 1 of meer cruciale voorspellende variabelen (CAUSALITEIT)*
 - Transparante verpakking tov niet-transparante, reclame met man tov met koppel, ...
- ★ *Meting van 1 of meer gedragsvariabelen*
 - Attitude tov advertentie, eetgedrag, ...
- ★ *Controle over randomisatie van andere potentieel voorspellende variabelen*

III. Ondervraging

= Een onderzoek of enquête aan de hand van een vragenlijst.

Kenmerken:

- Precieze vragen
- Kwantitatieve codering (antwoorden die kwantificeerbaar zijn)
- Forceert respondent in denkkader
- Omvang en representativiteit van steekproef zijn belangrijk

Opmaken van vragenlijst is ook een heel deel in de slides, maar alles is daar vrij logisch van.

2. Secundaire Bronnen

I. Interne secundaire data

In elke onderneming worden gegevens bijgehouden.

bv: verkopen per regio, per klant, klachten, ...

II. Externe secundaire data

A. Bronnen die voor iedereen toegankelijk zijn

Online websites met random info.

B. Bronnen die niet voor iedereen toegankelijk zijn

Tegen betaling, syndicated services (Nielsen, GFK)

Eerst doet men altijd een **kwaliteitscontrole** van de gegevens.

-> Wat was het doel van de studie?

-> Wie heeft de informatie verzameld?

-> Hoe is de informatie verkregen (was het een degelijke steekproef?)

-> Hoe consistent is de info met andere betrouwbare info?

Voorbeelden: Trade associations, Business press, Academic press, jaarrapporten, ...

Syndicated services: Op eigen initiatief en volgens een gestandaardiseerde procedure marktinformatie verzamelen en die tegen betaling ter beschikking stellen van de geïnteresseerde marktonderzoeker.

Deze informatie kan op 3 levels verzameld worden:



We zullen op de twee laatsten inzoomen.

Retailer panel

Een grote speler hier is **Nielsen**, dit bedrijf houdt scanner data bij van heel veel supermarkten (zoals Colruyt, Delhaize, etc..) -> Aldi en Lidl doen niet mee.

Ze houden dan scanner data bij, coupons & promoties en wat er op display wordt geplaatst.

Nielsen Data heeft 4 dimensies:

Market

- Where?
- Total Coverage
 - Coverage by Region
 - Shop type
 - Retailers

I. Market

Nielsen heeft België ingedeeld in verschillende regio's om de data te kunnen indelen in deze regio's.

De winkel zelf kunnen ook onderverdeeld worden in 'types'

- . F1 = Grote Hypermarkets (Colruyt, Carrefour, ...)
- . F2 = Gemiddelde grootte supermarkt ($\pm 400\text{m}^2$)
- . F3 = Kleinere winkels $< 400\text{m}^2$
- . F4 = 'Hard discounters' -> Aldi en Lidl

We zien wel dat de market shares van de kleinere winkels over de jaren sterk aan het dalen is.

Product

- Who?
- Category
 - Segment
 - Brand
 - SKU

II. Product

Bijhouden per product:

- . Segment waarin het hoort
- . Het merk
- . Het volume dat erin geraakt
- ...

Dit kan men bijhouden met behulp van de barcodes (unique identifiers)

Fact

- What?
- Sales – Value, Volume, Units
 - Price, Distribution

III. Fact

Volume Aantal eenheden, gram, aantal verkocht

Sales Aantal euro's verkocht

Prijs Prijs per eenheid, liter, gram, ...

Distributie IN hoeveel winkel beschikbaar (per soort product)

Period

- When?
- Weeks
 - 4-Weekly
 - Quarter,
 - Year (MAT or YTD)

IV. Period

Verkoop van producten zal op bepaalde moment populairder zijn.

Bv: Champagne tijdens de feesten

Ijsjes in de zomer

Is er een vaste periode tussen telkens een aankoop van dat product?

Consumentenpanel

Lange grote studies van een representatieve steekproef van consumenten.

Bepaalde informatie wordt herhaaldelijk gemeten bij dezelfde steekproefelementen.

GFK liet mensen thuis, bij het uitladen van hun boodschappen, ook scannen wat ze hadden gekocht om op niveau van consument patronen te herkennen en data daarvan bij te houden. Ze doen ook zelf analyses van deze data.

toev Nielsen geeft dit nog extra informatie:

- . Van waar komen hun kopers?

- . Zijn er loyale klanten

...

Ze houden ook bij welke producten de mensen kopen in combinatie met hoeveel ze tv kijken, welke reclame ze zien en dus ook de implicaties hiervan.

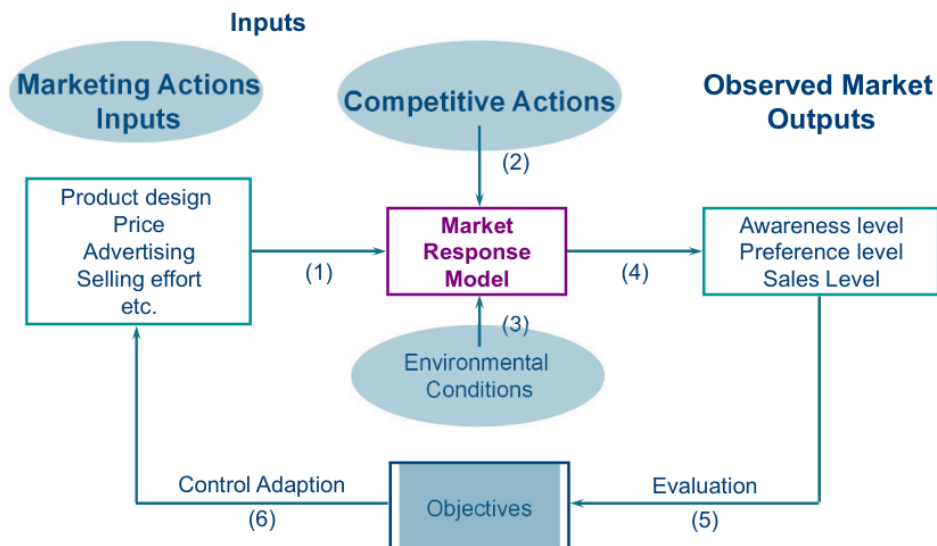
Nielsen en GFK zijn de bekendste spelers in deze markt, maar er bestaan er ook anderen.

Market Response Modellen

= De relatie tussen de marketing inputs (reclame, promoties, ...) en de invloed daarvan op de markt (verkoop, ...).

Dit zijn de bouwstenen van beslissingsmodellen, ze vereisen het expliciet maken van:

- ☐ **Inputs**
(on)controleerbaar
- ☐ **Response model**
Link tussen input en output
- ☐ **Objectieven**
Maatstaf om acties te kwantificeren/evalueren.



Er zijn verschillende soorten response modellen, die door volgende kenmerken gekarakteriseerd worden:

- **Aantal marketing variabelen** Meer variabelen zorgt voor grotere nauwkeurigheid.
- **Opnemen van competitie of niet**
- **Type van relatie** Lineair, kwadratisch, ...
- **Statisch of dynamisch** Kan er na het maken van het model nog input worden toegevoegd
- **Individuele of geaggregeerde response**
- **Niveau van vraag** (bv: Verkoop)

Belangrijke terminologie:

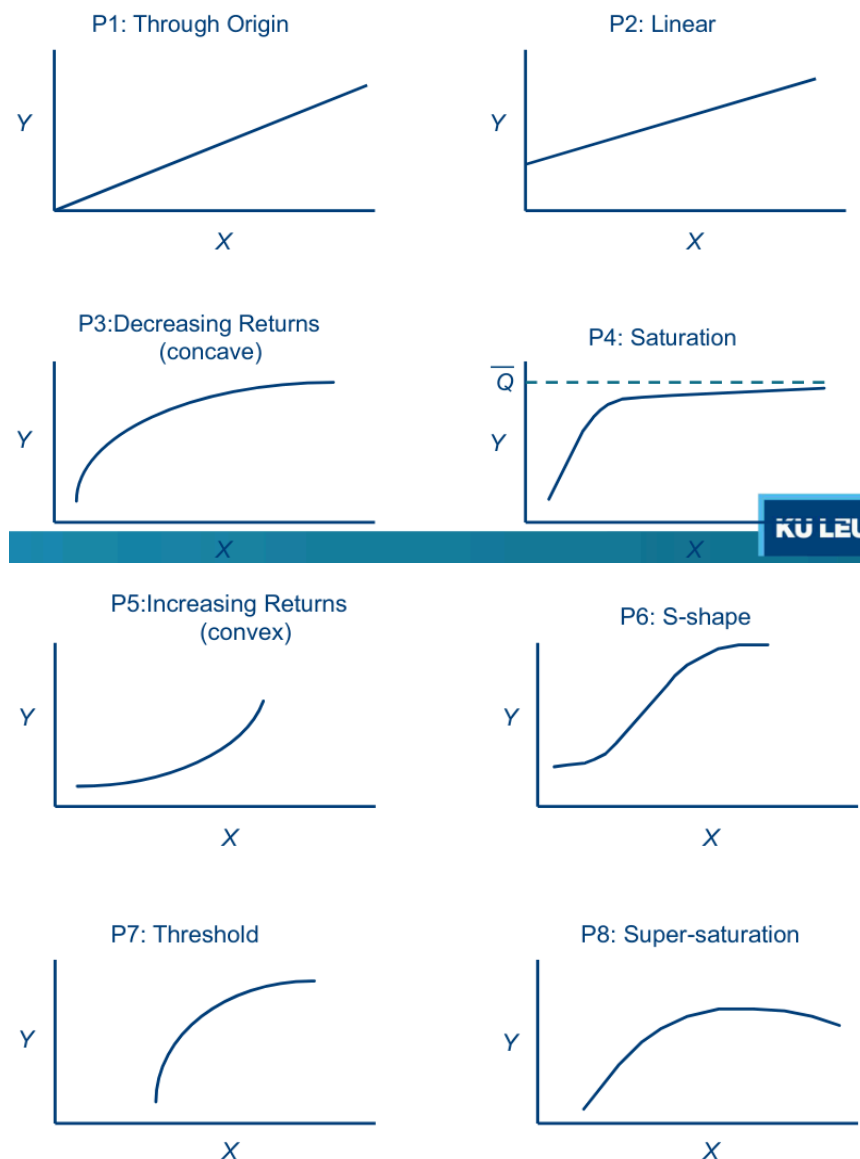
- **Afhankelijke vs onafhankelijke parameters**
- **Parameters** (constanten in de equatie)
- **Calibratie** (bepalen van gepaste waardes voor parameters)

Soorten Response Models:

1. Geaggregeerd response models
2. Gedisaggregeerde response models
3. Shared-experience models
4. Kwalitatieve response models

1. Geaggregeerde Response modellen

Fenomenen die kunnen voorkomen bij deze modellen



Modellen

I. Lineaire Model

$$Y = a + bX$$

Met X de reclame die je maakt voor dat product. Bij dit verband veronderstellen we

constante meeropbrengsten. Dit geeft wel onredelijk advies -> meer is beter, wat niet altijd het geval is.

II. Fractional Root Model

$$Y = a + bX^c$$

Lineair, stijgende of dalende opbrengsten (hangt af van C).

x

III. Exponential Model

$$Y = ae^{bx}; x > 0$$

Stijging of daling afhankelijk van b.

IV. Modified Exponential Model

$$Y = a(1 - e^{-bx}) + c$$

Dalende meeropbrengsten en saturatie (a+c) met ondergrens (c)

V. ADBUDG Model

$$Y = b + (a-b) \frac{X^c}{d + X^c}$$

S-shaped voor $c > 1$

concaaf van vorm voor $0 < c < 1$,

b = ondergrens, a = bovengrens (saturatie).

Vaak gebruikt om response modellen/ verkoop effort te modeleren, Toepasbaar voor judgemental calibration.

Calibratie

= Het toewijzen van goede waarden voor de parameters in het model.

Goede benadering bekomen van hoe Y varieert met waarden van X.

Least squares: Minimaliseert de som van de gekwadratiseerde verschillen tussen elk geobserveerde Y en de geschatte waarde voor Y.

I. ADBUDG Model: Subjectieve calibratie

Op basis van 5 vragen aan de verkoper:

- Wat is het huidige niveau van reclame en verkoop?
- Wat is de verkoop als de hoeveelheid advertenties = 0?
- Wat is de verkoop als de hoeveelheid advertenties daalt met 50%?
- Wat is de verkoop als de hoeveelheid advertenties stijgt met 50%?
- Wat is de verkoop als de hoeveelheid advertenties = ∞ ?

Hiermee kan men dan de parameters bepalen en het model opstellen.

II. Objectieven

We proberen hiermee marketing advies te evalueren en de prestatie van het bedrijf te verbeteren.

Voor korte-termijn winst: Eenheidsmarge * Verkochte hoeveelheid - relevante kosten

Uitbreiding

De nieuwe variabelen kunnen opgeteld (additief model) of vermenigvuldigd worden (multiplicatief model) of beide (multiplicatief + additief model).

note: Het multiplicatieve model wordt zeer vaak gebruikt in marketing

I. Meerdere variabelen

bv. tussen reclame en prijs (zijn de X) en Y is dan de verkoop

1. Additieve model: $Y = a \cdot f(X_1) + b \cdot g(X_2)$

-> De variabelen X moeten onafhankelijk zijn, er is dus geen verband tussen reclame en prijs

2. Multiplicatieve model: $Y = a \cdot f(X_1) \cdot g(X_2)$

-> De variabelen gaan afhankelijk zijn, er is een interactie tussen beide, dat de impact van de ene een versterkend of juist verzwakkend effect kan hebben op het andere.

3. Multi + Add model: $Y = a \cdot f(X_1) + b \cdot g(X_2) + c \cdot f(X_1) \cdot g(X_2)$

Dit is wel vooral aangeraden als X_1 en X_2 niet te sterk gecorreleerd zijn, dus als hun correlatie < 0.8

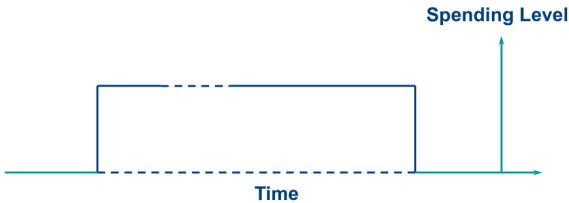
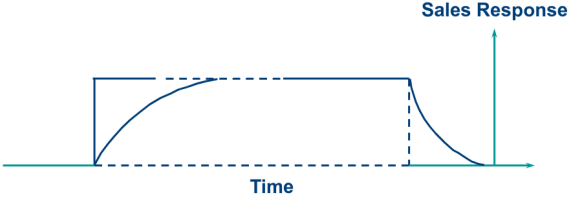
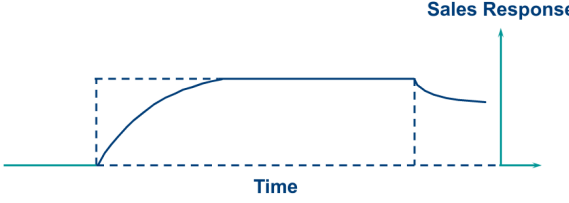
Als je een model hebt met meerdere X, verklarende variabelen, dan moet de multicollineariteit kleiner zijn dan 0.8. mogen niet te sterk gecorreleerd zijn!

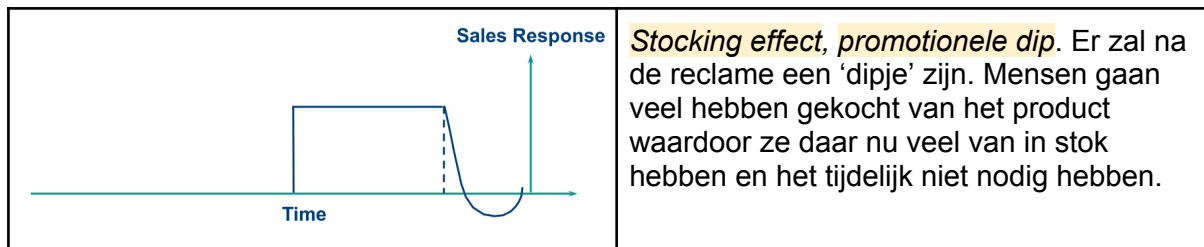
Een ander multiplicatief model: $Y = aX_1^b X_2^c$

Waar b en c elasticiteiten zijn. Dit model wordt zeer vaak gebruikt in marketing.

|-> Elasticiteit = Met hoeveel % verandert Y als X met 1% is veranderd.

II. Dynamische Effecten

 Spending Level Time	Op deze grafiek zien we duidelijk het Carry-over effect: De invloed van marketingactie in deze periode op verkoop, marktaandeel, ... in latere periodes.
 Sales Response Time	Hier ziet men duidelijk delayed response en customer holdout. Dit betekent dat mensen niet onmiddellijk op een actie zullen reageren, of dat er na de actie nog reactie zal blijven nazinderen.
 Sales Response Time	Hysteresis effect: Er zal saturatie optreden na de reclame omdat de kopers het product aantrekkelijk blijven vinden. Doel van elke marketeer om dit te bereiken, er is een blijvend effect van de marketingactie.
 Sales Response Time	Hier ziet men het new trier en wear out effect, reclame zal voor een hogere piek zorgen maar dan uiteindelijk terug zakken naar een normaler niveau.



Het respons model hiervoor is eigenlijk een lineaire model dat daarbovenop rekening houdt met de acties van het verleden ($t-1$). Dus op die manier ook alle andere acties die die actie voorging:

$$Y_t = a_0 + \underset{\substack{\uparrow \\ \text{current} \\ \text{effect}}}{a_1 X_t} + \underset{\substack{\uparrow \\ \text{carry-over} \\ \text{effect}}}{\lambda Y_{t-1}} + \text{error}$$

III. Marktaandeel (-> Competitie)

Y -> Brand sales models

V -> Product Class sales models (Hoe groot is die categorie)

M -> Market share models

$Y = M \times V$ (= Aandeel in de markt van dat product)

Market share attraction model: $M_i = \frac{A_i}{A_1 + A_2 + \dots + A_n}$

2. Gedisaggregeerde response modellen (Eerder informatief)

Probability of choice: Een individu's probabiteit/kans dat die merk 1 kiest.

$$P_{i1} = \frac{e^{A_1}}{\sum_j e^{A_j}}$$

$$A_j = \sum_k w_k b_{ijk}$$

$A_j = \sum_k w_k b_{ijk}$ is de attractiveness van een product j voor individu i.

- b_{ijk} is de mening van individu i over het merk j op basis van een kenmerk (bv. kwaliteit) k.
- w_k bepaalt hoe belangrijk een kenmerk k is.

Segmentatie, Targeting & Positionering

1. Segmentatie

= Het proces om mensen in te delen in groepen waarbij:

- Mensen in eenzelfde groep gelijkaardig zijn
- Mensen in een verschillende groep echt verschillend zijn.

De segmenten zijn enkel nuttig als elk segment verschillend zal reageren op marketing projecten.

Stappenplan voor STP:

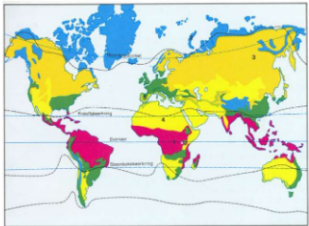
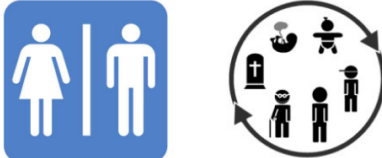
- Segmentatie. Identificeer groepen van klanten (segmenten).
 - Stap 1: Marktsegmentatie
 - Stap 2: Segmenten beschrijven
- Targeting. Kies een segment waaraan je wil verkopen.
 - Stap 3: Evalueren attractieve (goede) segmenten
 - Stap 4: Selecteren van doel segmenten & het alloceren van middelen
 - Stap 5: Vinden van de consumenten in de doelgroep
- Positioning. Plaats product in de gedachten van de klanten
 - Stap 6: Identificeer mogelijk positioneringconcepten voor elke doelgroep
 - Stap 7: Selecteer, ontwikkel en communiceer het gekozen concept

Fase 1: Marktsegmentatie

- > We gaan de markt al opdelen in verschillende groepen

1. Beweegreden: Waarom segmenteren we deze markt?
2. Selecteer een set van segmentatie variabelen.
3. Selecteer een mathematische & statistische procedure om consumenten samen te nemen in segmenten.
4. Groepeer consumenten in een vastgelegd aantal segmenten.

<i>Geografische segmentatie</i>	Stedelijk vs landelijk, afhankelijk van landen, dorpen, continenten, ...
---------------------------------	--

	vb: Lays heeft in Nederland een joppiesmaak en hier in België de bickysmaak. Grootte van producten kan ook afhangen van land tot land, bv de cola flessen zijn in Amerika standaard 2 liter.
Demografische segmentatie 	Man/vrouw/x, levensfase van die persoon, inkomen, beroepen, opleiding, gezinsgrootte, ... vb: Lego is gemaakt voor leeftijd x - y De deo's voor mannen en vrouwen apart, ...
Psychografische segmentatie 	Levensstijl, persoonlijkheid & waarden. vb: Ecologisch wasmiddel voor mensen die belang hechten aan het klimaat, biologische voeding, ...
Gedragsmatige segmentatie 	<u>Beste segmentatie om je op te baseren!</u> - <i>Gelegenheid</i> het vliegtuig nemen voor het werk of juist voor op vakantie -> verschillende noden - <i>Voordelen</i> Mensen willen op een reis verschillende dingen, .. - <i>Gebruikersstatus</i> Niet gebruikers, voormalige, mogelijke, nieuwe, bestaande - <i>Mate van gebruik</i> Light, medium en heavy users - <i>Loyaliteit</i> Loyale, semi-loyale, tijdelijke loyale en switchers. - <i>Houding</i> Enthusiast, positief, onverschillig, negatief en vijandig.

|-> Hiervoor moet men eigenlijk te weten komen wat er in de hoofden van de mensen leeft.

Fase 2: Marktsegmenten beschrijven

Hier zijn twee belangrijke begrippen te onderscheiden:

-> Segmentatievariabelen (Bases)

Kenmerken die aangeven waarom segmenten verschillen. Deze zijn dus uniek binnen segmenten en belangrijk om deze segmenten te kunnen definiëren.

-> Beschrijvende variabelen (Descriptors)

Kenmerken die helpen om segmenten te vinden en te bereiken, ook wel onafhankelijke variabelen genoemd.

Segmenten zullen het beste bereikt worden met variabelen die makkelijk te observeren zijn. bv (heel cru) geslacht, bruin of blond haar, woonplaats, ...

2. Targeting

Met deze segmenten kan je twee dingen doen:

Diversificatie strategie

= keuze van verschillende segmenten maken

- + De moeilijkheden van 1 segment zo opvangen door succes te hebben in andere markten.

Focus strategie

= Op 1 marktsegment focussen

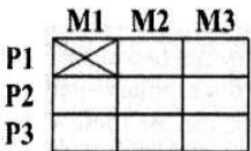
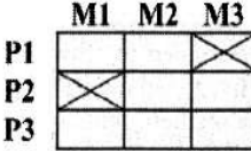
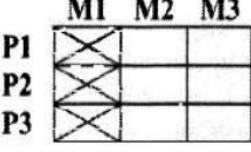
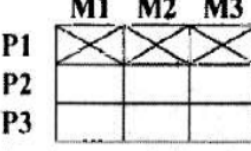
- + meer gerichte positionering, voornamelijk in zeer competitieve markten.

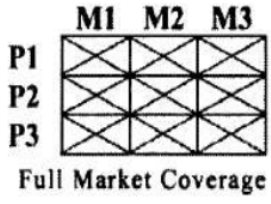
Fase 3: Evalueren attractiviteit segmenten

Om een segment te kiezen kan je de attractiviteit van een segment overwegen:

- Hoe groot is de vraag in elk segment en groeit dat segment nog?
- Is er een grote concurrentie in dat segment?
- Is er een gat in de markt in dat segment?
- Kan je het segment afschermen voor anderen?
- Past het segment bij mijn al bestaand bedrijf (vb: Colgate dat eten begon te maken)

Fase 4: Selecteer doelsegmenten & alloceren van middelen

 Single Segment	1 segment in 1 markt (<i>Focus strategie</i>) -> 1 product zonder veel variatie. vb: Ferrari, zij gaan niet veel variëren in hun auto's en zijn vooral gericht op rijke mensen die sportwagens willen.
 Multi Segment	Segmenten verspreid over een paar markten (<i>Diversificatie</i>) vb: Mcdo, deze zijn begonnen met enkel hamburgers en frietjes. Later kwam daar ontbijt bij (vrij duidelijk voor een ander doelpubliek) en ze hebben tijdelijk pizza's gebruikt, wat niet succesvol was.
 Market Specialization	Een bedrijf gaat voor 1 markt alles voorzien Dupont: voorziet explosieven voor alle soorten mijnen die er zijn, er zijn bedrijven die allerhande soorten auto's voorzien (wat dus allemaal in diezelfde markt zit)
 Product Specialization	Een bedrijf gaat voor een bepaald segment alles voorzien vb Boeing maakt vliegtuigen voor alle vliegtuig bedrijven, allerhande soorten, ...

	Mja, heel de markt inpalmen op alle segmenten. vb Harry Potter: voor alle leeftijden en ook met films, boeken, speeltjes, etc.
---	--

Fase 5: Vinden van consumenten i/d doelgroep

Ofwel zal de consument zelf moeten zoeken naar wat hen aanstaat, ofwel zal deze een vragenlijst kunnen invullen om dan advies te krijgen over welk product voor hen perfect is.

Als men onderzoek gaat doen, kan men de data van dat onderzoek allemaal verzamelen en analyseren mbv een *clusteranalyse*.

SEGMENTATIE ANALYSE

Segmentatie kan met behulp van verschillende methodes

- Factor analyse (data reductie)
- Cluster analyse (Segmenten vormen)
- Discriminant analyse (Segmenten beschrijven)

Clusteranalyse

Maatstaf definiëren om de mate van gelijkenis tussen klanten te meten op basis van hun noden.

- Groepeer klanten met gelijkaardige noden.
 - Hierarchische clustering (Ward's methode)
 - Niet hierarchische clustering (K-Means)
- Aantal segmenten kiezen op basis van strategische criteria, eigen oordeel en cijfers.
- Een profiel maken van de noden van de geselecteerde segmenten = cluster profiling wie zijn die segmenten?

Maak een profiel van de noden van de geselecteerde segmenten.

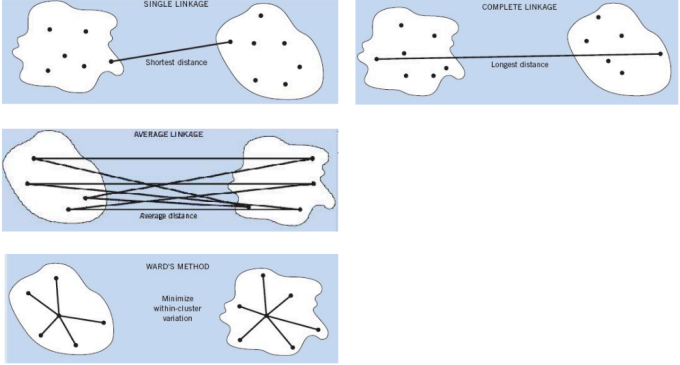
> Maatstaf voor gelijkenis

<i>Matching Coëfficiënt</i>	Hoe vaak is hetzelfde antwoord gegeven tussen de proefpersonen
<i>Afstandsmaatstaf</i>	Letterlijke afstand op een geplote weergave. Die kan met de <i>Euclidische afstand</i> $\sqrt{(x_{1i} - x_{1j})^2 + (x_{2i} - x_{2j})^2 + \dots + (x_{ni} - x_{nj})^2}$ <i>Absolute afstand</i> $ x_{1i} - x_{1j} + x_{2i} - x_{2j} + \dots + x_{ni} - x_{nj} $

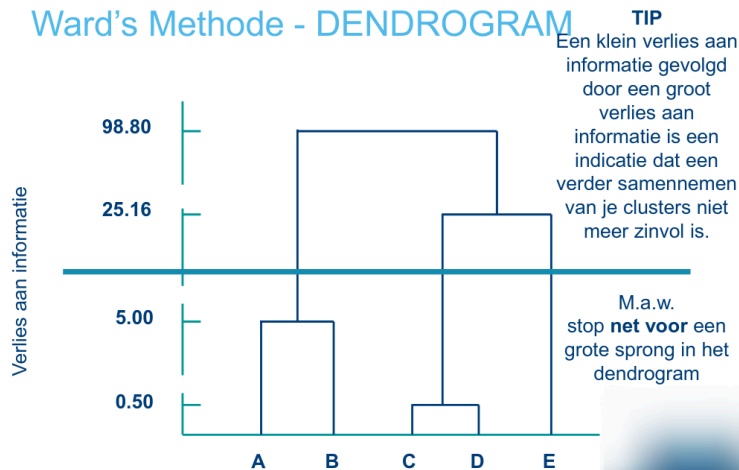
$$x \mapsto \frac{x - \bar{x}}{\hat{\sigma}_x}$$

Bij het bepalen van die afstanden, moet rekening gehouden worden dat ze op dezelfde schaal staan, dus dan zal men moeten *standaardiseren*.

> Cluster methodes

<i>Hiërarchische methode</i>	Elke respondent in een apart segment, dan de meest gelijkende samenvoegen en zo proces verder doen tot er 1 groot segment is met iedereen. Hier zal men een stopcriterium definiëren om op tijd het juiste aantal segmenten te hebben. afstand afhankelijk van welke methode je gebruikt:  <u><i>Ward's methode</i></u> bovenop deze methodes gaan we proberen zorgen voor een minimaal verlies aan informatie. -> som van gekwantificeerde afwijking van elke afwijking met zijn clustergemiddelde Voor het <i>aantal</i> clusters zal men een compromis zoeken tussen A. Veel en dus meer homogene clusters B. minder clusters en beter datareductie -> dit is dus een <i>exploratieve methode</i> Je voert analyse best uit met verschillende clusteraantallen.
<i>Niet-Hiërarchische methode</i>	x segmenten blijven door elkaar yeeten tot je het juiste aantal optimale clusters hebt. Voor het aantal clusters zal men kijken naar het dendrogram (zie onder tabel) Bij elk resultaat zal men moeten kijken of de clusters voldoende verschillend zijn, of er clusters samengenomen moeten worden of uiteengetrokken. Ook zal men zien of je logische namen kan plakken op de segmenten.

Ward's Methode - DENDROGRAM



➤ Cluster profiling

Hiermee wil men nakijken of een gekozen groep clusters wel de juiste keuze was, dit doen we met behulp van ANOVA.

ANOVA zal twee soorten aan variatie meten voor data en zal deze dan vergelijken.

Variation BETWEEN groups

Voor elke data waarde kijk je naar het verschil tussen dat groepsgemiddelde en het algemene gemiddelde.

Variation WITHIN groups

Voor elke waarde kijk je naar het verschil van de waarde en het gemiddelde van de groep.

ANOVA is dan uiteindelijk de ratio van de variatie tussen groepen tegenover de variatie binnen de groepen:

$$F = \frac{\text{Between}}{\text{Within}} = \frac{MSG}{MSE}$$

Een grote F zal tegen de hypothese zijn, aangezien er dan meer verschil is tussen groepen dan in de groepen zelf.

-> kies dus non-gerelateerde beschrijvende variabelen voor je segmenten en hou rekening met de aantrekkelijkheid en sterktes van eigen bedrijf in die variabelen.

3. Positionering

= Strategieën die een bedrijf ontwikkelt om hun product te differentiëren in de hoofden van zijn doelgroep. Het product moet een onderscheidende, belangrijke houdbare positie bekleden in de hoofden van de consument.

Verschil met differentiatie: Het creëren van tastbare of ontastbare verschillen op 1 of meerdere attributen met je belangrijkste concurrenten. (materieel je product verbeteren)

Fase 6: Identificeer mogelijke positioneringsconcepten voor elke doelgroep

De producent probeert dus dat de consumenten het product op een bepaalde manier bekijken: *uniek*, *superieur* of gewoon *gelijkaardig maar goedkoper*.

Voor uniek is bijvoorbeeld het hoogste hotel in een land hebben een goed voorbeeld, Coca cola en Pepsi proberen al jaren elkaar te ondermijnen en de superieure frisdrank te lijken en vele huismerken zijn uiteindelijk hetzelfde als gewone merken maar gewoon goedkoper.

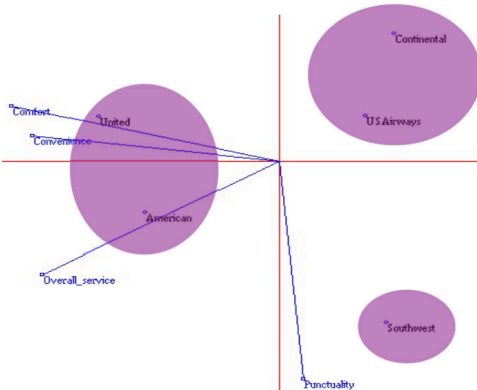
Je kan inspelen op:

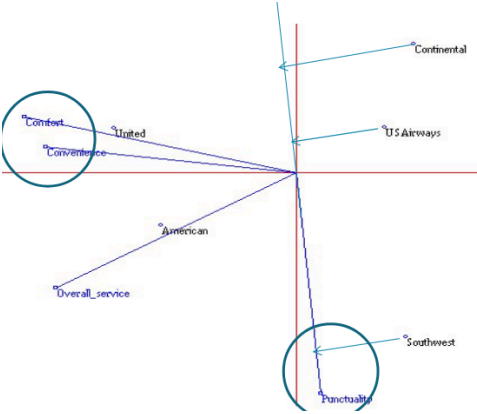
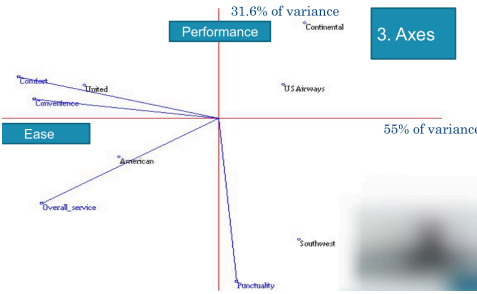
- **Levensstijl**
inzetten op levensstijl van de mensen, zoals ecologisch zijn of juist vrij ongezond eten.
- **Attribuut**
Inzetten op een attribuut van het product, zoals bijvoorbeeld de zachtheid van toiletpapier, ..
- **Voordeel**
Inzetten op een voordeel, zoals dat 'Seven-up' de dorst lest.
- **Competitie**
Coca cola en Pepsi zitten onder constante competitie tegenover elkaar, waardoor mensen ook een kant kiezen en ermee bezig zijn.
- **Tijd**
Aangeven wanneer je het product kan gebruiken, wanneer het handig zou zijn. Zoals de Royco soepjes tijdens de pauzes op het werk.

Managers gebruiken een techniek die toelaat om differentiatie & positionering strategieën te ontwikkelen door een competitieve structuur van hun markt, zoals gepercipieerd door hun klanten te visualiseren.

-> Perceptuele map

Grafische weergave waar concurrerende alternatieven zijn geplott in een euclidische ruimte, deze hebben 3 belangrijke kenmerken:

<i>Distances</i> 	Hoe meer afstand er is tussen bepaalde producten/producenten, hoe minder gelijkaardig ze zijn.
<i>Vectors</i>	Hoe kleiner de hoek tussen twee vectoren (van attributen), hoe groter de correlatie tussen die attributen. Dit geldt ook voor hoeken van 180°, deze zijn namelijk perfect tegenovergesteld gecorreleerd. De grootste hoek is in dit geval dus 90°. Wanneer een vector lang is, betekent dit dat deze perfect wordt opgenomen door de map (tegenover

	de assen). Korte vectoren lijden aan informatieverlies.
Axes 	De assen zullen benoemd worden op basis van de attributen die daar in de buurt komen. We projecteren al deze attributen op twee assen, dus er zal informatieverlies optreden. In het algemeen zal dit geen groot probleem zijn.

Je kan in deze mappen soms gaten zien waar sommige bedrijven nog niet op zijn ingegaan. Deze gaten kunnen er zijn omdat dat letterlijk een gat is in de markt en de mensen dit willen hebben, ofwel omdat dit niet nodig is in de markt. Daarom kan men ook eens de meningen van een steekproef op de map loslaten. Op deze manier kan je de gaten in de markt accurater bepalen en onderzoeken of ze daadwerkelijk winstgevend zouden zijn.

Toepassingen van perceptuele mappen:

- I. **Nieuw product beslissingen**
De gaten in de markt ontdekken, testen met wat jij op de markt zou brengen.
- II. **Check van competitieve structuur en positionering**
Nakijken waar in de markt je je bevindt en wat je concurrenten hun positie is.
- III. **Identificeer tegen wie je concurreert**
Wie zijn mijn concurrenten, voor wie moet ik opletten met mijn product.
- IV. **Imago/reputatie studies**
Wat is mijn imago volgens de wereld, is mijn imago veranderd sinds kort, ..

Factoranalyse

= Techniek die systematisch onderliggende patronen/dimensies en onderlinge relaties tussen attributen opspoort. Op die manier kan men de factoren voor de assen bepalen.

Elke factor is een combinatie van lineaire attributen:
$$F_j = a_{j1}x_1 + a_{j2}x_2 + \dots + a_{jn}x_n$$

Idealiter gebruiken we 2 dimensies/ 2 factoren:

- **Factor 1:** Neemt zoveel mogelijk informatie in de attributen X's op.
- **Factor 2:** Orthogonaal tot Factor 1 en neemt zoveel mogelijk van de resterende informatie in de attributen X's op.

Elke observatie in de dataset kan bij benadering ook weergegeven worden door een lineaire

combinatie van factoren $x_{kj} \approx z_{k1}f_{1j} + z_{k2}f_{2j} + \dots + z_{kr}f_{rj}$.

Fase 7: Selecteer, ontwikkel en communiceer het gekozen concept

-> wordt niet meer echt gezegd, miss nieuw product enzo voor bezien.

Toevoegen van voorkeuren op de perceptuele map

Ideal point method	Introduceer een ideaal product/bedrijf alsof het een extra stimulus is vanuit je consumenten - Je kan analyseren hoe gelijkaardig jouw product hiermee is - Je kan ratings van attributen van het ideale product combineren met andere producten. => nieuw product op map
Vector method	Introduceer voorkeuren als additionele variabelen in de attributen data - Je kan deze voorkeuren evalueren samen met de andere producten, welk attribuut beïnvloedt de voorkeur het meest, welke producten liggen het dichtst bij deze voorkeuren. - Je kan de voorkeuren ook apart bekijken, dan kan je je segment targetten. => nieuwe vector op map

Preference maps:

= Wat als respondenten gelijkaardige percepties, maar zeer verschillende voorkeuren vertonen:

- Ontwikkel een perceptuele map voor concurrerende alternatieven
- Bereid de map uit met de voorkeuren voor elke respondent
 - Identificeer doelsegmenten
 - Bereken een verwachte marktaandeel

Het marktaandeel kan volgens 2 manieren berekend worden:

1. First choice of maximum utility rule

Consument kiest product dat die het meest prefereert

	Voorkeur Respondent 1	Voorkeur Respondent 2	Voorkeur Respondent 3
A	3	<u>5</u>	4
B	<u>4</u>	3	<u>5</u>
C	3	2	1

$$\text{Marktaandeel A} = (0+1+0) / 3 = 33.33\%$$

$$\text{Marktaandeel B} = (1+0+1) / 3 = 66.67\%$$

$$\text{Marktaandeel C} = (0+0+0) / 3 = 0\%$$

Dit is van toepassing bij duurdere producten of producten waar lang over nagedacht moet worden voordat het wordt gekocht

II. *Share of preference rule*

Consument kiest elk product in hun consideratie set in verhouding tot de relatieve voorkeur voor dat product in vergelijking met de andere producten.

	Voorkeur Respondent 1	Voorkeur Respondent 2	Voorkeur Respondent 3
A	3	5	4
B	4	3	5
C	3	2	1

A	3 / (3+4+3)	5 / (5+3+2)	4 / (4+5+1)
B	4 / 10	3 / 10	5 / 10
C	3 / 10	2 / 10	1 / 10

$$\text{Marktaandeel A} = (3+5+4) / 3 = 40\%$$

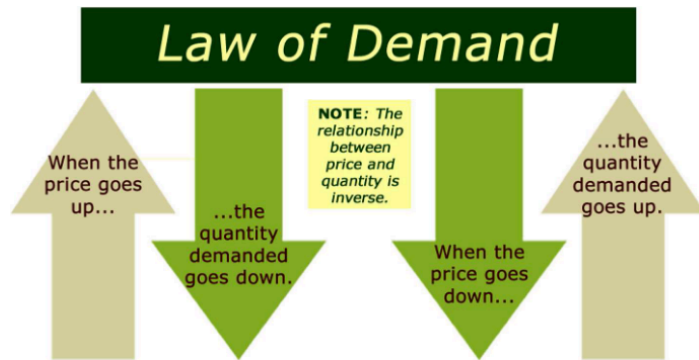
$$\text{Marktaandeel B} = (4+3+5) / 3 = 40\%$$

$$\text{Marktaandeel C} = (3+2+1) / 3 = 20\%$$

Dit is voor dingen zoals bier of eten, etc.

Pricing

Nodige theorie om dit deel te snappen



-> Dit zou het geval zijn in een perfecte wereld waar mensen zich volledig out of context baseren op de prijs op dat moment van dit product.

In de echte wereld gaan mensen soms een daling in prijs niet merken, of die daling gebruiken om te weten dat er een nieuw product komt om dat dan te kopen.

Een andere uitzondering zijn producten die in waarde stijgen met hun prijs, waarbij demand soms kan stijgen met de prijs.

- **Quantity demanded:** $f(\text{Price})$
- **Profit:** $\text{Quantity}(\text{Price}) * [\text{Price} - \text{Unit Cost}]$
- **Maximize profit** -> where marginal cost = Marginal revenue

$$\left[\frac{1}{E} + 1 \right] P = MC \Rightarrow P = \left(\frac{1}{\frac{1}{E} + 1} \right) MC$$

- Prijs elasticiteit

= ratio van het percentage verandering in demand tegenover het percentage verandering in de prijs.

= het ratio waarmee de vraag stijgt of daalt in verhouding met de prijs.

Elasticity Ratio	Type of Elasticity	Description
$E = \infty$	Perfectly elastic	Any very small change in price results in a very large change in quantity demanded.
$1 < E < \infty$	Relatively elastic	Small changes in price cause large changes in quantity demanded.
$E = 1$	Unit elastic	Any change in price is matched by an equal change in quantity.
$0 < E < 1$	Relatively inelastic	Large changes in price cause small changes in quantity demanded.
$E = 0$	Perfectly inelastic	Quantity demanded does not change when price is changed.

Prijs bepalen in praktijk

Cost-Oriented pricing	De kosten gaan de prijs het meest beïnvloeden. Voordelen: Dit is populair aangezien het gemakkelijk is om in te schatten, gemakkelijk te rechtvaardigen en het simplificeert de prijsbepaling. Nadelen: Je houdt geen rekening met de buitenwereld, misschien had je wel meer kunnen vragen.. Dit is dus een slechte bepaling.
Competitor-Oriented pricing	Going-rate pricing: Bedrijven proberen hun prijs op het gemiddelde van de huidige markt te houden. vb: Goed voor dingen zoals olie, dit is een homogeen product dat

	niet sterk zal variëren onder bedrijven, dus is het constant houden van de algemene prijzen belangrijk. Competitive bidding: Een bedrijf is in competitie met een onbekend nummer aan producenten, zonder informatie over hun prijzen. vb: Bijvoorbeeld wanneer een gemeente een nieuwe straat wilt aanleggen en aan verschillende aannemers anoniem een prijs vraagt. Dan zal de beste prijs winnen maar de aannemers weten niet van elkaar wat de prijs is.
Demand-Oriented pricing / Value-Based pricing	= Proberen een hogere prijs te vragen wanneer demand hoger is en een lagere prijs wanneer deze lager is. Dit is onafhankelijk van de productiekosten. Value-in-use: De prijs waarbij de mogelijke koper onverschillig is tegenover het blijven kopen van het huidige product of het veranderen naar het nieuwe product. : <i>true economic value</i>.

VIU = Productie kost + Winst marge + Consumenten surplus

Be sure to include all costs when doing a VIU calculation.

- (Annual) Purchase cost
- Fabrication cost
- Finishing cost
- Inventory cost
- Maintenance/Service cost
- Scrap adjustment
- Level-of-requirement adjustment
- Changeover cost
- Risk premium

Zoeken van value-in-use in een voorbeeld:



Value-in-Use

A chemical plant uses 200 O-rings to seal valves carrying corrosive materials. Those O-rings cost \$5.00 each and must be changed during regular maintenance every two months.

A new product has 2 x the corrosive resisting power. The value-in use of the material might be:

Annual cost of incumbent

= 200 * 6 * \$5/O-ring
= \$6,000
= 200 * 3 * VIU

VIU = \$10



Value-in-Use

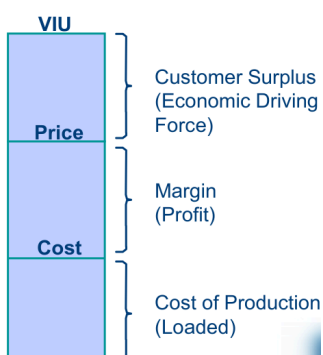
Suppose the new material allows a longer time between shutdowns—4 months vs. 2 months, and the cost of a shutdown is \$5,000.

200 * 6 * 5 + 5,000 * 6 = 200 * 3 * VIU + 5,000 * 3

↓ ↓ ↓ ↓
 Equipment Shutdown Equipment Shutdown
 Cost Cost Cost Cost

Incumbent New

VIU = \$35



De producent zal altijd onder de VIU moeten gaan aangezien de klanten anders niet zullen veranderen van product. Dit is een soort bovengrens voor de prijs van een product.

Anderzijds zal zou de koper het liefst een zo laag mogelijke prijs zien en dus zo dicht mogelijk bij de kosten dat het nam om dat product te maken. Dit is een soort ondergrens voor het bepalen van de prijs.

De **True Economic Value (TEV)** is de totale waarde die een product (bv. een machine) biedt. Deze is vaak hoger dan de verwachte waarde en dus de VIU.

Perceived value: De waarde die de kopers aan de producten toekennen. Deze is spijtig genoeg altijd lager dan de TEV:

- Kopers beseffen de voordelen van het product niet
- Ze beseffen de voordelen, maar zijn sceptisch
- Ze beseffen volledig de voordelen, maar snappen niet dat dit belangrijk kan zijn.

⇒ Solution: Verbeter het level en de kwaliteit van de marketing campagnes gericht op de klant.

Dit kan men dus verbeteren door meer reclame te maken om de koper te "overtuigen/informereren van de voordelen. De prijs moet dan ook wel logisch zijn:



Als de producenten zouden weten hoeveel waarde een koper aan hun product geeft, kunnen deze hun prijs aanpassen om mensen die er meer aan willen geven, meer te laten betalen tegenover mensen die er niet veel voor willen geven. Dit noemt men **Price customization**.

Price customization

Het is moeilijk exact, en volledig eerlijk vanuit de koper, te bepalen hoeveel een persoon wilt uitgeven aan een product. Dit gebeurt wel al zonder het expliciet te vragen.

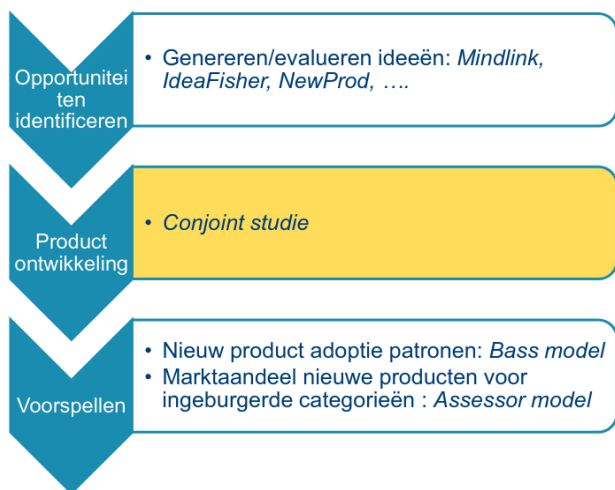
Arbitrage: prijs kan afhangen per land. Dit kan ervoor zorgen dat retailers hun spullen uit een ander land halen om het in hun land duurder te verkopen.

Price customization kan ook Illegaal zijn. Solden voor 1 bepaald geslacht is bijvoorbeeld bij wet verboden. Kortingen per leeftijd mogen dan wel (zoals voor studenten of 65+ers). Deze kortingen kunnen ook oneerlijk aangevoeld worden door mede kopers omdat zij dan geen korting krijgen.

Online kan dit wel gemakkelijker en stiekemer aangezien je per persoon een andere prijs kan instellen zonder dat de koper dat ziet. Deze moet dan al op meerdere accounts gaan kijken of er verschillen zouden zijn.

Prijsdiscriminatie is het aanpassen van de prijs aan wat en segment bereid is te betalen. Dit is een manier om aan revenue management te doen.

Non-linear Pricing/Quantity discounts	<i>korting bij grote hoeveelheden:</i> <i>Two-part tariff:</i> Lidmaatschap in clubs (de fitness) Bestaat uit een vast bedrag om de x tijd te betalen + constante variabele kosten <i>Block tariff:</i> Een korting op hoeveelheid (vaak in de Colruyt met de extra kaart)
Temporal Price discounts	<i>Time-of-day pricing</i> <i>Time when purchased</i> <i>Day of the week pricing</i> <i>Seasonal pricing</i> ... <i>Revenue management:</i> De kunst en wetenschap om het juiste product aan de juiste klant op het juiste moment met de juiste prijs te verkopen. vb: vliegtuigtickets verkopen aan mensen die dat vroeg kopen aan een lage prijs of business mensen die twee dagen op voorhand zijn.
Geografische prijs discriminatie	<i>In sommige landen is men bereid een hogere prijs te betalen voor een product.</i> <i>Als resultaat kan de winst gemaximaliseerd worden door hogere prijs voor eenzelfde product te vragen</i>



Nieuw Product

-> Een nieuw product kan een zeer grote investering zijn en kan gemakkelijk falen. Hier zijn dus ook modellen rond gemaakt.

We gaan ons in dit hoofdstuk vooral focussen op de Conjoint analyse, om te zien wat wat de voorkeuren van de consumenten zijn.

Voor deze analyse gaan we producten identificeren aan de hand van hun attributen

en bepalen hoe belangrijk elk attribuut is.

Conjoint analyse

(leer dit best met het voorbeeld van de pizza's, dat maakt het duidelijker)

= Een methode om de structuur van consumentenvoorkeuren te verstaan en te incorporeren in het nieuw-product ontwikkelingsproces.

Het laat toe om te evalueren hoe consumenten een **trade-off** maken tussen verschillende product-attributen en -levels.

Output:

- Numerieke waarde voor het relatieve belang dat consumenten geven aan attributen in een product categorie
- Numerieke waarde voor het relatieve belang dat consumenten hechten aan elk kenmerk/level binnen een attribuut.
- Identificatie van product designs die marktaandeel of andere indices maximaliseren.

Je kan met deze analyse een vergelijking maken tussen attributen die niets met elkaar te maken hebben (appelen met peren vergelijken).

Waarom werkt het?

- Respondenten moeten een afweging maken, net als in de echte wereld.
- We ontnemen respondenten de kans om te zeggen dat alle attributen even belangrijk zijn
- Wanneer respondenten gedwongen worden een afweging te maken leren we wat ze echt belangrijk vinden.

Stage 1- Designing the conjoint study

1.1 Welke attributen

Deze kunnen gekozen worden op basis van:

- Literaire werken
- Diepgaande interviews met management
- Market analysis
- Interviews met consumenten

We gaan maximaal 6 attributen voor een product kiezen, om juist genoeg informatie bij te kunnen houden.

1.2 Welke levels in de attributen

een rule of thumb die we raadplegen is 2-5 levels binnen een attribuut.

Deze levels zouden aan volgende kenmerken moeten voldoen:

- Ze representeren de laagste tot de hoogste waarde in de markt.
- #levels per attribuut moet ongeveer gelijk zijn tussen de attributen.

1.3 Bundels van producten maken

We kunnen alle mogelijk combinaties maken van de levels van de attributen, maar dan verliezen we het feit dat we mensen duwen om een keuze te maken. Er zal dus een selectie moeten gebeuren.

Fractional Factorial Design: Een subset nemen uit alle productcombinaties.

Bij het bepalen van de subset moet men rekening houden met *Orthogonality*.

= Voor elk level van een attribuut, zullen de levels van elk ander attribuut met dezelfde ratio moeten voorkomen. ..

Stage 2 - Obtaining data from a sample of respondents

2.1 Design a data-collection procedure

- Pairwise evaluations of product bundles
Personen telkens tussen 2 bundels laten kiezen.
- Rank-ordering product bundles
Een hele lijst geven aan bundels en mensen ze laten ranken
- evaluating products on a rating scale
Elk product een rating geven tussen 0 en 100

2.2 Select a computation method for obtaining part-worth functions.

$$R_{ij} = \sum_{k=1}^K \sum_{m=1}^{M_k} a_{ikm} x_{jkm} + \varepsilon_{ij}$$

- j: a particular product/concept included in the study
R_{ij}: the rating provided by respondent i for product j
a_{ikm}: part-worth associated with the mth level (m=1,2,..., M_k) of the kth attribute
M_k: number of levels of attribute k
K: number of attributes
x_{jkm}: 1 if the mth level of the kth attribute is present in product j, 0 otherwise
ε_{ij}: error term

Waarbij a ook vaak gegeven is. Na deze computatie kan men voor een bepaald persoon (met hun voorkeuren gegeven) de utility berekenen. Dit door elke keer de part-worth waarvoor zij kiezen te nemen en deze allemaal op te tellen, of in een moeilijke formule:

$$U(p) = \sum_{k=1}^K \sum_{m=1}^{M_k} \hat{a}_{ikm} x_{pkm}$$

- p: a particular product/concept of interest
U(P): the utility associated with product P
a_{ikm}: estimated part-worth associated with the mth level (m=1,2,..., M_k) of the kth attribute
M_k: number of levels of attribute k
K: number of attributes
x_{pkm}: 1 if the mth level of the kth attribute is present in product P, 0 otherwise

Stage 3 - Evaluating product design options

3.1 Segment customers based on their part-worth functions

_____ spreekt wel voor zich

3.2 & 3.3 Design market simulations and select choice rule

_____ Definieer de set waaruit kopers zullen kiezen, mbv de choice rules:

<i>Maximum utility rule/ First choice rule</i>	Elke koper kiest het product dat hen het meeste utility geeft. De market share wordt dan gegeven door: $MS(P_i) = \sum_{k=1}^K \frac{\text{Consumers_who_prefer_product_the_most}}{K}$ K is het aantal mensen dat heeft deelgenomen aan de studie.
<i>Share of preference rule</i>	Elke koper kiest een product met een kans proportioneel met de utility ervan vergeleken met de totale utility van alle producten. $p_{ij} = \frac{u_{ij}}{\sum_{j=1}^I u_{ij}} \quad MS_j = \frac{\sum_{i=1}^I p_{ij}}{I}$

Profit potential of a new product

Market share * Incremental margin over base product

Assess cannibalization potential of new product

Een voor en na analyse doen

- Market share zonder de nieuwe producten.
- Market share met de nieuwe producten.

Situaties waarbij conjoint analyse handig kan zijn:

- Ontwikkelen van een nieuw product dat de meeste waarde creëert voor de consument.
- Voorspellen van verkoop/ marktaandeel van alternatieve producten
- Identificeren van segment die een gelijkaardig voorkeuren hebben voor een product bundel.
- Identificeren van het beste product voor een bepaalde doelgroep
- Het nagaan van de impact van verschillende prijszetting en service strategieën.

Voor-nadelen om voorkeuren te weergeven

Voordelen

Respondenten worden gedwongen een trade-off te doen, gelijkaardig aan een echte aankoopssituatie

Nadelen

Klassieke limitaties van een vragenlijst

- is steekproef representatief
- is wat de proefpersonen zeggen effectief wat ze doen
- Mensen kunnen verveeld worden
- ...

Voor-nadelen om een nieuw product te maken

Voordelen

Respondent moet een trade-off maken (veel accuratere voorspellingen)

Simuleert de hele markt, zonder de kost van een echte introductie.

Winkellocatie

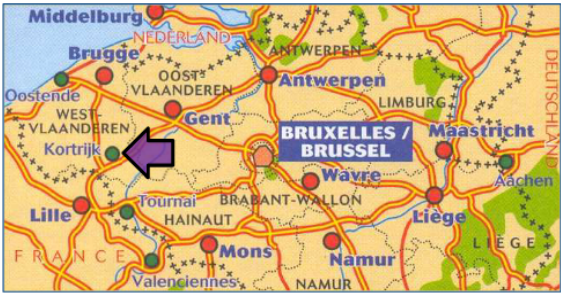
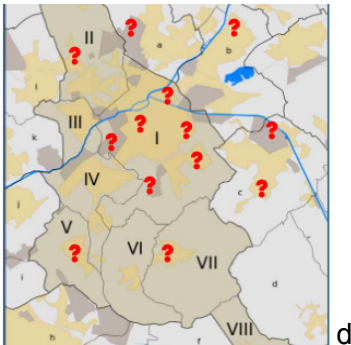
Dit is de belangrijkste variabele van een winkel. Waarom?:

- Nadelen slechte locatie zijn moeilijk te overwinnen
- Bepaalt wie je aantrekt
- Lange termijn commitment & niet gemakkelijk te veranderen
- Gepaard met grote financiële investeringen

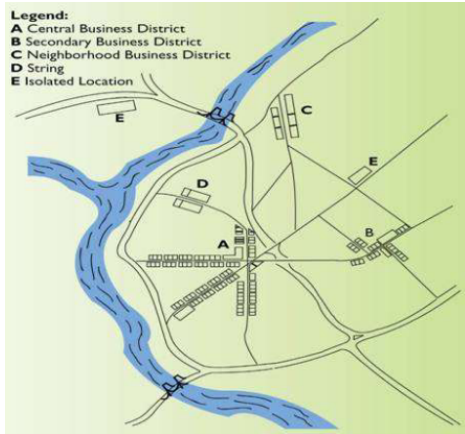
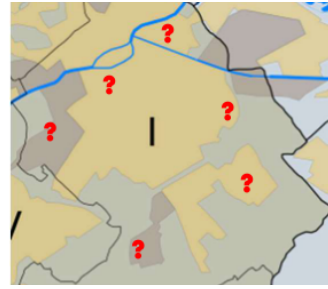
Bij een locatie moet je ook met vele criteria rekening houden:

- I. Populatiegrootte en kenmerken
- II. Concurrentie
- III. Toegankelijkheid
- IV. Zichtbaarheid
- V. Parkeermogelijkheden
- VI. Types nabijgelegen winkels
- VII. Huurkosten
- VIII. Wettelijke restricties
- IX. Lengte huurovereenkomst
- X. ...

Het beslissen van de locatie kan in 3 stappen gebeuren:

Selectie van (lokale) markt(en)	
Gebiedsanalyse	

locatie-evaluatie



Er zijn ook verschillende winkellocaties:

E = geïsoleerde winkel

Ongepland business district

A = Centraal BD

B = Secundair BD

C = buurt BD

D = string

Gepland business district

Selectie van locaties

Dit kan op twee manieren:

Op basis van ervaring?

96% van retailers

Op basis van checklists:

Populatie

Welk kenmerk heeft de bevolking in die regio?

grootte, leeftijdsgroepen, gezinssamenstelling, beroep, levensstijl, ...

-> huidige en potentiële koopkracht

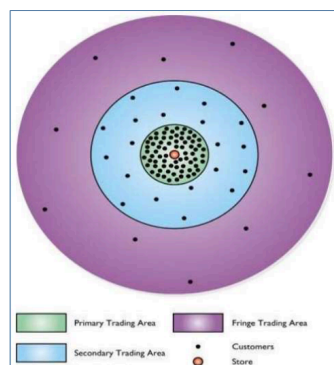
Bereikbaarheid

Is er veel passage aan het pand, is het bereikbaar met de metro, is er een bushalte dichtbij of een metrostation, is het veilig voor de fiets en de voetgangers.

Is er een makkelijke manier voor de leveranciers om de goederen te brengen.

Trading area

= geografisch gebied dat de kopers bevat van een bepaald bedrijf of winkel voor specifieke goederen. -> vorm hoeft niet perse bolvormig te zijn.



- **Primair verzorgingsgebied**
= 50-80% van alle klanten
- **Secundair verzorgingsgebied**
= 15-25% van alle klanten
- **Randgebied (= fringe)**
= alle overige klanten



Grootte en vorm van dit gebied zijn afhankelijk van enkele kenmerken:

- Grootte van de winkel
- Nabijgelegen winkels
- Vervoersnetwerk
- Bevolkingsdichtheid
- Fysieke, sociale en politieke barrières
- locatie van de concurrentie
- Marketing strategie: Hoe uniek het assortiment is, lage prijzen, superieur, ...

Methodes om dit gebied vast te stellen:

1. *Afleiden van reeds bestaande winkels*

In kaart gebracht via loyaliteitskaarten, vragenlijsten, postcodes, ... -> kan gecombineerd worden met bevolkingsstatistiek.

2. *Definiëren a.d.h.v. verplaatsingstijden*

Indicatie voor 3 niveaus verzorgingsgebied.

Concurrentie

Mensen willen bij een winkel-uitje, graag in 1 keer al hun boodschappen doen. In het geval van supermarkten en kledingwinkels, is het dus voordelig je langs je concurrenten te zetten. Verschillende types winkels kunnen ook indirect concurrenten zijn.

Anchor stores: Dit zijn winkels die alom bekend zijn en waar mensen van heel ver naartoe zullen gaan. Als je winkel daar langs ligt, zal je een grote kans hebben dat een deel van deze klanten ook bij jou langskomen.

Winkels kunnen ook complementair zijn, wat natuurlijk ook voor attractie zorgt.

Index of retail saturation

= index met de potentiële verkoop per m² in een verkoopgebied.

$$\frac{\begin{array}{c} \# \text{ customers for given} \\ \text{product line} \end{array} \times \begin{array}{c} \text{Avg. retail expenditures} \\ \text{for given product line} \\ \text{per customer} \end{array}}{\begin{array}{c} \text{Square footage of selling space allocated to a given product} \\ \text{line within competing retail facilities} \\ \text{in a market} \end{array}}$$

Hier wordt wel geen rekening gehouden met de concurrenten of met je eigen positionering.

Kosten

- Aankoop/huurprijs
- Renovatie/bouw/aanpassing interieur/...
- Bouwvoorschriften: beperkingen?
- Onderhoudskosten: lift, veiligheid,...
- Belastingen
- Leveringskosten, ...

De huurprijs van een groot pand op een goede plek zal al veel geld vragen en je winkel moet dus winstgevend genoeg zijn hiervoor.

Beoordeling van locaties

-> voorspellen van omzet op verschillende locaties:

Analoge methode

Op basis van reeds bestaande winkels:

1. **Identificatie van gelijkaardige locaties**
2. **Kwantificatie van relevantste winkelmerken**
3. **Extrapolatie naar eventuele winkels op overwogen locaties**

Spatial interaction (Gravity)

Klantenstromen voor een winkel zijn gerelateerd aan:

<i>Aantrekkelijkheid van de winkel</i>	<i>Toegang Barrières</i>
- Vloeroppervlakte - Winkelimago - Lokale concurrentie	- Afstand - Parking - Veiligheid - Omgeving

Dit kan via Huff's Gravity model:

Het model van Huff bepaalt de kans dat een klant naar de winkel zal komen gebaseerd op de vloeroppervlakte van alle relevante winkels in de buurt en hun afstand tot de klant. Met de gedachten dat:

- Een klant minder snel ver gaat
- Een klant zal sneller naar een grotere winkel gaan

-> De kans dat een gegeven klant zal winkelen bij een specifieke winkel of winkelcentrum wordt groter naarmate de winkel(centrum) groter naarmate de afstand of reistijd voor de klant daalt.

De kans wordt berekend met de ratio van het winkeloppervlak t.o.v. de reistijd van iedere winkel te berekenen. Door dit ratio te vergelijken met dat van de andere winkels wordt de kans gevonden.

P_{ij}	kans dat persoon (i) locatie (j) bezoekt
S_j	door locatie (j) toebedeeld vloeroppervlakte voor gegeven productcategorie
T_{ij}	reistijd voor persoon (i) naar locatie (j)
λ	parameter voor effect reistijd op verschillende trip types
$\sum_{j=1}^n$	som over alle locaties (n)

$$P_{ij} = \frac{\frac{S_j}{(T_{ij})^\lambda}}{\sum_{j=1}^n \frac{S_j}{(T_{ij})^\lambda}}$$

De constante lambda zal afhangen waarvoor mensen naar deze winkel gaan. Voor melk zal je niet ver willen reizen, waardoor je lambda zeer laag zal zijn. voor een nieuwe keuken zal je verplaatsing al wat groter mogen zijn, dus dan is je lambda hoger.

Regressie

Voorspellen omzet op nieuwe locatie o.b.v. statistische link tussen omzet en kenmerken van

bestaande locatie, specificatie: $Omzet_j = \beta_0 + \beta_1 * X_{1,j} + ... + \beta_k * X_{k,j}$

Het is belangrijk te bepalen welke variabelen we meenemen.