

De modellering van data uit het verleden voor een advies over de classificatie van de SKU.

Roel van den Hoogen

S0090522

Technische Bedrijfskunde

Universiteit Twente

Voorwoord

Het afronden van deze opdracht is voor mij het einde van mijn studie en daarom wilde ik dit graag bij een bedrijf doen om zo alle kennis die ik heb opgedaan tijdens mijn studie tot uiting te brengen in een echt project. Daarom was ik heel blij toen ik werd geaccepteerd om voor de Wim Bosman Group een case aan te pakken, die voldoende diepte gaf en toch nog direct invloed had op de resultaten van het bedrijf. Vandaar dat ik me met veel plezier heb ingezet, om deze opdracht zo goed mogelijk uit te voeren.

De eerste week heb ik mogen meedraaien op de werkvloer in Ans20 en ik ben dankbaar voor die tijd die iedereen heeft besteed om mij bekend te maken met de procedures, die gebruikt worden binnen het warehouse.

Vervolgens heb ik ongeveer 2 maanden op het kantoor gewerkt, om de case uit te werken tot een werkend programma. Hierbij heb ik mezelf meester gemaakt van Visual Basic en Excel om zo een bijna autonoom lopend programma af te leveren. Sander Memelink Teamleader en Iris Timmermans Operations Manager hebben mij hier bijgestaan in alle vragen of problemen, die de case met zich meebracht. Verder heb ik elke week overleg gehad met Thomas Bijl Branch Manager Logistic Services om het project op koers te houden en het beste resultaat te garanderen.

Door de kennis die ik heb opgedaan tijdens mijn studie heb ik zelfstandig aan de opdracht gewerkt en de veel problemen zelf kunnen oplossen. Daarnaast heb ik de punten waar ik vast liep met behulp van de tips van Dr. Peter Schuur van de faculteit Management en Bestuur (MB) kunnen oplossen. Op grond van zijn begeleiding heb ik het verslag tot een samenhangend geheel kunnen vormen.

Ik wil graag iedereen bedanken die mij hebben geholpen deze opdracht af te ronden.

Als laatste wil ik mijn ouders bedanken, die mij hebben gemotiveerd, om een stage binnen een bedrijf te zoeken. Zonder hun motivatie weet ik niet of ik de studie wel zou hebben afgerond.

Management summary

Voor deze bacheloropdracht heb ik stage gelopen binnen de Wim Bosman Group.

Het warehouse heeft een suboptimale indeling van de locaties van de SKU en ze hebben mij gevraagd om te kijken of verbetering hiervoor mogelijk was. Er wordt op dit moment gebruik gemaakt van een indeling in zones. Het doel was een programma in Excel te schrijven om snel en efficiënt een toewijzing van de SKU aan een zone te geven.

De probleemstelling bij dit project was:

Hoe kan de data uit het verleden worden gemodelleerd tot een advies over de classificatie van de SKU?

Om deze probleemstelling aan te pakken zijn de drie onderzoeksvragen aangepakt:

Welke variabelen zijn van belang voor het bepalen van een goede classificatie?

Op grond van literatuuronderzoek en na overleg is gekozen voor de variabele: Aantal order lines (per maand).

Hoe kan er rekening gehouden worden met de vervuiling van de data zoals reclames/ acties voor een SKU zodat geen verkeerde conclusies worden getrokken?

Dit wordt ondervangen door te werken met periodes van een maand en het gebruik van de variabele 'aantal order lines'.

1. Het gebruik van maandelijkse periodes houdt in dat acties die een week duren minder opvallend in de data terugkomen.
2. Het 'aantal order lines' is een variabele die minder invloed ondervindt van acties en reclames door het feit dat het aantal casepicks hierin niet terugkomt. Het aantal casepicks zal omhoog gaan met een reclame (meer dozen verkocht), maar de order lines zullen hier geen directe invloed door hebben aangezien er voldoende cases op een pallet liggen. (over het algemeen gaan er 5 tot 100 order lines uit een pallet). Zelfs een order met waarbij 1 SKU heel veel cases bevat zal bijna nooit van meer dan 2 pallets moeten picken aangezien deze anders uit de bulk gehaald worden.

Door deze twee keuzes zal de variantie in de vraag weinig reactie tonen op reclames en acties.

Naast de vervuiling van acties en reclame was er nog een meer directe vervuiling in de vorm van SKU's, die wel in de data voorkomen, maar geen onderdeel worden van een van de 4 zones. Om deze vervuiling tegen te gaan wordt is het mogelijk deze uit de data te filteren voordat een planning gemaakt wordt. Dit wordt met behulp van macro's in Excel, automatisch gedaan.

Welke methodes van terug/vooruit kijken zijn geschikt voor de indeling in slow/fast?

Deze onderzoeksvraag was lastig direct uit de literatuur te halen, omdat er veel verschil in de gebruikte methodieken is. De twee overheersende methodieken zijn corston's method en exponential smoothing. Beide methodieken worden veel gebruikt in het bedrijfsleven en worden vaak geleverd als planning tools in een ERP programma. Er is veel data voor nodig om een goede inschatting te krijgen.

In deze case variëren de SKU's jaarlijks sterk en de SKU's zijn ook seizoensgebonden. Dit maakt het lastig ze in te schatten met deze twee methodieken.

Tevens is de meeste literatuur met betrekking tot voorspellingen niet gericht op het classificeren van SKU, maar de literatuur is gericht op voorspellen van exacte voorraad.

Vanuit deze optiek is gekozen om een variant op de regression analyse te gebruiken voor de berekening van een gewogen gemiddelde. Er wordt een dataset (aantal maanden) bepaald als input. De maanden zullen een gewicht toegedeeld krijgen op grond van het feit hoe ver de maanden van de target maand af liggen.

Inhoudsopgave

VOORWOORD	I
MANAGEMENT SUMMARY	II
INHOUDSOPGAVE	IV
HOOFDSTUK 1: INLEIDING	1
HOOFDSTUK 2: PROBLEEMSTELLING	2
ONDERZOEKSVRAGEN	2
VERWACHTE RESULTATEN	3
DOELSTELLINGEN	3
PLAN VAN AANPAK	4
TIJDSPLANNING	4
HOOFDSTUK 3: LITERATUUR	6
TOEWIJZING VAN PRODUCTLOCATIES	6
FORECASTING STRATEGIE	7
BEPALING AANTAL CLASSES	8
BEREKENING BESPARING	8
SAMENTVATTING LITERATUUR	11
HOOFDSTUK 4 : KEUZES	12
STORAGE ASSIGNMENT METHOD	12
VARIABELE VOOR TOEWIJZING LOCATIES	12
FORECASTING STRATEGIE	13
AANTAL CLASSES	13
HOOFDSTUK 5: IMPLEMENTATIE	13
HOOFDSTUK 7: CONCLUSIE	14
AANBEVELINGEN	15
HOOFDSTUK 8: GEBRUIKTE LITERATUUR	17

Hoofdstuk 1: Inleiding

In het kader van deze bacheloropdracht heb ik bij Wim Bosman Group stage gelopen om de de classificatie van de SKU aan te pakken.

Wim Bosman Group is onderdeel van “Mainfreight” en zij bieden global supply chain logistiek in meer dan 200 locaties in Europa, USA, Azië, New Zeeland, Australië en Zuid Amerika.

De Wim Bosman Group is opgericht in 1963 en heeft 1600 medewerkers in Europa. Wim Bosman Group heeft verschillende warehouses over heel Europa en verzorgt hiermee de logistiek en het distributie voor verschillende klanten. Naast de opslag en het transport levert de Wim Bosman Group ook Value Added Logistics in de vorm Co-Packing en andere services.

Voor deze opdracht is gekeken naar een specifieke operatie met Fast Moving Consumer Goods binnen het warehouse in 's-Heerenberg en daarom zal eerst iets worden verteld over de methodieken die gebruikt worden binnen het warehouse.

De producten (SKU = stock keeping unit) worden aangeleverd op volle (bulk) pallets en deze worden binnen het warehouse op A (grond) niveau gelegd of op het B of hoger niveau. Op A (grond) niveau liggen aangebroken pallets en als deze op raken wordt een pallet vanuit de bulk in de correcte zone geplaatst.

Bij het order picken is er onderscheid tussen bulk (volle pallet) en casepick (cases en eatches) orders. Bulk orders worden uit de bulk zones gepickt. Casepick orders worden op een orderpick wagen op grond niveau gepickt volgens de route die is gepland in het warehouse management systeem (WMS).

Als de volledige order is verzameld, of de pallet vol is, wordt er terug gereden naar de staging waar de pallet wordt afgeleverd, gecontroleerd en ingepakt. Vervolgens worden de pallets in de vrachtwagen gereden voor transport.

In dit project wordt gekeken naar het toewijzen van een zone aan elke van de SKU in het warehouse. Om dit te bereiken, zal in Excel een macro geschreven worden, die automatisch de belangrijke gegevens uit de ruwe data filtert en vervolgens op grond van deze gegevens een advies geeft, met betrekking tot de zone van elke SKU.

Daarvoor heb ik onderzoek verricht naar het bepalen van de productlocaties binnen een warehouse. De sub gebieden hierin zijn methodieken van toewijzing, forecasting methodes, bepaling van het aantal classes en de kostenberekening.

Hoofdstuk 2: Probleemstelling.

De producten (SKU = stock keeping unit) voor de operatie liggen samen in een warehouse voor klant B. De SKUs zijn op dit moment ingedeeld in productcategorieën met de daarvoor bestemde zones. Deze categorieën bestaan uit fast/slow, heavy/light en de uitzonderingen. Deze SKU hebben echter een sterke variantie in de vraag, omdat ze vaak seizoensgebonden zijn. Op dit moment is er geen tool om de categorieën snel en eenvoudig te bepalen en moet dit met de hand worden gedaan. Dit kost echter te veel tijd en wordt slechts sporadisch gedaan, met als gevolg dat de seizoensgebonden producten niet op de juiste locatie liggen. Het is mogelijk om aan de hand van de data uit het verleden een relatief goede inschatting van de vraag voor de toekomst te maken.

Het doel bij deze opdracht is het maken van een tool die met behulp modelering van data uit het verleden een advies kan geven voor de classificatie van de SKU in slow of fast.

Hiervoor moet de data worden geprepareerd. Dit wordt gedaan door te filteren op uitzonderingen en de resultaten worden om te zetten in een draaitabel. Nadat deze data is omgezet naar een overzichtelijke tabel, wordt met behulp van toekennen van een gewicht aan elke periode, van de resultaten een gewogen gemiddelde berekend. Dit gemiddelde zal leiden tot een advies over de classificatie Slow/Fast. Vanuit dit advies kan het management de toewijzingen van SKU's aan nieuwe zones veranderen in het systeem.

De probleemstelling die hierbij hoort is als volgt:

Hoe kan de data uit het verleden worden gemodelleerd tot een advies over de classificatie van de SKU?

Onderzoeksvragen

De probleemstelling kan worden opgedeeld in verschillende deelvragen die het makkelijker maken om de probleemstelling te beantwoorden.

Om deze deelvragen te vinden zal gebruik worden gemaakt van een probleemkluwen, om de vragen te vinden, die van belang zijn om de probleemstelling aan te pakken.



De focus ligt op de rechter boom van deze probleemkluwen.

Dit houdt in dat begonnen wordt met de volgende drie deelvragen:

- Welke variabelen zijn van belang voor het bepalen van een goede classificatie?
Om deze vraag goed te beantwoorden is dit in 3 stappen gedaan.
 1. Welke storage assignment methode is ideaal voor deze Case?
 2. Bij gebruik van class based storage: Hoeveel classes is ideaal?
 3. Welke variabele is van belang voor de toewijzing van de SKU?
- Hoe kan er rekening gehouden worden met de vervuiling van de data zoals reclames, acties voor een SKU, zodat geen verkeerde conclusies worden getrokken?
Door de verschillende onregelmatigheden wordt de data sterk beïnvloed en aangezien de resultaten uit het verleden beperkt zijn zal het van belang zijn om door middel van verschillende smoothing methodes op goede gebalanceerde inschatting te komen voor de uiteindelijke indeling in slow/fast
- Welke methodes van terug/vooruit kijken zijn geschikt voor de indeling in slow/fast?
Veel producten zijn seizoensgebonden en zullen dus een enkele keren per jaar van zone veranderen en het is belangrijk dit op tijd te doen om de producten in de juiste zone te hebben liggen tegen de tijd dat het seizoen begint
- Verder is er nog gekeken hoe er kan berekend worden wat de nieuwe methode bespaard per maand.

Verwachte resultaten

De verwachting is dat aan het eind van het project er een Excel tool beschikbaar is dat vanuit data een classificatie slow/fast kan maken per SKU. De berekeningen, die gebruikt worden in dit bestand zullen ondersteund worden met behulp van literatuur en metingen op de vloer in het warehouse.

De classificatie slow/fast geeft vervolgens een productcategorie. Zo kan er met minimale input elke maand de zone (categorie) worden bepaald per SKU.

Doelstellingen

De doelstellingen die ik voor mezelf heb gesteld zijn als volgt:

- Het leren toepassen van de opgedane kennis tijdens mijn studie in een praktijksituatie en hierin het bedrijf helpen om een probleem op te lossen. (inzichtaspect/vaardigheidsaspect)
- Het oplossen van een praktisch probleem en hierbij wetenschappelijke literatuur toepassen voor het verbeteren van de oplossing en vergaren van nieuwe kennis (vaardigheidsaspect)
- Leren werken met collega's die niet aan hetzelfde project werken, maar wel inzichten en kennis van zaken hebben en zo behulpzaam kunnen zijn (vaardigheidsaspect)
- Initiatief tonen met het vinden van een eigen bacheloropdracht (houdingsaspect)
- Inzicht krijgen in de organisatie om zo de problemen beter in kaart te brengen en op te lossen waar dit binnen de scope ligt (inzichtaspect)

Plan van aanpak

1. Er zal begonnen worden met het vergaren van kennis over de producten en de werkmethodes binnen het warehouse. Dit wordt voornamelijk gedaan door een week mee te draaien op de vloer, om zo ervaring op te doen met het werk dat hier gedaan wordt, om zo de kritische punten binnen het proces te kunnen onderscheiden.
Het proces is op de vloer is als volgt:
Na de ervaring op de vloer ben ik aan de slag gegaan met de analyse van het probleem en het plan van aanpak. Voor de analyse is er bekeken hoe de huidige indeling in categorieën is bepaald, want dit geeft informatie over de eisen en de variabelen die van belang zijn. Tevens is in overleg bepaald welke filters er over de data gehaald moeten worden, voordat deze geschikt is om conclusies uit te trekken.
2. Vervolgens zullen de probleemstelling, de probleemkluwen en onderzoeksvragen uiteengezet worden.
3. De huidige situatie wordt kort beschreven.
4. Er zal een literatuuronderzoek worden gedaan voor verschillende gebieden.
 - a. Methodes van toewijzing van productlocaties.
 - b. Hoeveel classes optimaal zijn voor deze case.
 - c. Key variabele van toewijzing.
 - d. Methodes van smoothing en planning toepasbaar zijn op deze case.
 - e. De achtergrond van de savingsberekening wordt kort uiteengezet.
5. De gegevens van de literatuur worden toegepast op de operatie en er wordt een selectie gemaakt uit de beschikbare opties.
6. Implementatie van de case houdt in dat er 2 Excel programma's moeten worden geprogrammeerd:
 - a. De eerste Excel tool heeft als doel het filteren van de data. De overgebleven data wordt omgezet in een draaitabel met daarin de overgebleven resultaten.
 - b. Nadat bekend is welke methodes optimaal zijn voor deze case zullen deze geprogrammeerd worden in een tweede Excel macro bestand. Dit bestand zal aan de hand van input van de gebruikers data uit de draaitabellen halen. Deze data wordt vervolgens omgezet in een berekening per SKU, voor een gewogen gemiddelde variabele. Op grond van dit gewogen gemiddelde, zullen alle SKU's een categorie krijgen. De gebruiker kan vervolgens deze data invoeren in het systeem en hiermee elke maand de optimale setup voor het warehouse instellen.
7. Na de implementatie zal de Tool getest worden op verschillende aspecten en op grond van de resultaten zal een conclusie en verscheidene aanbevelingen worden gegeven.

Tijdsplanning

Week	Planning
21	Meelopen op de vloer, bekend worden met werkmethodes en bedrijf
22	Maken plan van aanpak, Excel gebruiken om de data te filteren en draaitabel te maken
23	Literatuur onderzoek bepalen variabelen, meten data variabelen
24	Programmeren nieuwe data en begin programmeren voor combinatie datasheets
25	Literatuur onderzoek smoothen, averagen, seizoen curves, data verleden eventuele start programmeren
26	Programmeren, bijwerken handleiding en uitwerking data

27	Programmeren, bijwerken handleiding en uitwerking data
28	Schrijven verslag, bijlage en controle code
29	Afronding, presentatie gebruikers,
30	Afronding

Hoofdstuk 3: Literatuur

In dit hoofdstuk wordt de literatuur uiteengezet hierin wordt onderscheid gemaakt in 5 onderdelen.

- Toewijzing van productlocaties
- Forecasting strategie.
- Variabele voor toewijzing locaties.
- Bepaling aantal classes
- Savingsberekening

Toewijzing van productlocaties

De volgende keywords zijn gebruikt tijdens het zoeken: warehouse, class-based, storage, batching, order (line), location, assignment, continued, adaptation, seasonal.

Storage assignment methods

Om een overzicht te maken van de mogelijke assignment methodes zullen de belangrijkste stukken uit verschillende artikelen worden opgesomd. Op sommige punten zal kort toelichting gegeven worden.

Random storage. ^{(1) (2) (3) (4)}

Random plaatsen van een SKU.

Closest open location ⁽¹²⁾

SKU plaatsen op de eerste locatie die vrij is binnen het warehouse vanaf het I/O punt gezien. Het is aangetoond dat dit vergelijkbaar is met random storage op gebied van besparing.

Dedicated storage ^{(1) (2)}

Alle SKU een vaste locatie geven op grond van een criteria.

Class based storage ^{(3) (4)}

De SKU indelen in 1, 2, 3, etc classes. 1 class voor alle SKU is random storage. Zoveel classes als dat er SKU zijn is dedicated storage.

Family based storage. ⁽¹⁾

Dit houdt in dat een familie van producten in dezelfde zone liggen. Dit kan op gewicht zijn, maar ook op vergelijkbare producten.

Variabelen toewijzing locaties in een warenhuis.

De volgende variabelen zijn in de literatuur gebruikt ter bepaling van de class/locatie van een SKU.

1. Family, similar SKU and similar size. ^{(1) (4) (9)}

Family is lastig te bepalen en zal bij groot verloop snel veranderen. Er is wel veel potentie voor winst om familie te gebruiken als variabele. Als orders slechts een aisle in moeten is dit erg efficiënt.

2. Throughput / Largest demand of volume per period per SKU ^{(1) (4) (7) (8) (9)}

De locatie van de SKU wordt gebaseerd op de throughput van elke SKU. Dit houdt in dat SKU met een hogere throughput dicht bij het IO punt liggen. Een nadeel van deze methode is dat er sprake is van 'Traffic intensity blocking' en dat er drukke aisles zijn bij de SKU met een hoge throughput en rustige aisles bij de SKU met een lage throughput. Als er weinig ruimte is in de gangen kan dit leiden tot verstopping en extra wachttijden tijdens het picken. Er is een significant vermindering van picking time geresulteerd uit de kortere afgelegde afstand.

3. Demand frequency ^{(3) (4) (9) (10) (11) (12)}

De locatie van de SKU wordt gebaseerd op de frequentie dat er een doos gepickt wordt. (vergelijkbaar met throughput) Dit is een methode die nuttig is als er verschillende aantallen casepicks zijn bij de orders. Hiermee wordt bespaard op de rijafstand en maakt de tijd die het kost om cases op de pallet te laden niet mee genomen.

4. Duration-of-stay of a single unit (pallet, case, each) / Normal duration of stay (DOS). ^{(4) (5) (8) (9)}

Door uit te rekenen hoe lang een SKU op een locatie ligt kan een indicatie gegeven worden hoe vaak deze moet worden bijgevuld en hiermee kan de ideale locatie bepaald worden. Deze methode is optimaal voor bulk picking. Deze methode vereist een grote hoeveelheid berekeningen en data die direct beschikbaar is.

5. Cube-per-order index (COI) is gedefinieerd als de ratio van de vereiste ruimte van een SKU tegenover de populariteit (het aantal opslag/afhaal verzoeken). ^{(2) (4) (6) (9)}

Deze methode wordt voornamelijk gebruikt voor bulk picking.

6. Pick probability (the higher the pick probability the closer the item to I/O), passed orders

Forecasting strategie.

Voor het bepalen van de strategie met betrekking tot de weging van data uit het verleden is in de literatuur gekeken naar vergelijkbare problemen.

Hierbij zijn de volgende keywords gebruikt bij het zoeken: moving average, exponential, smoothing, inventory, production, planning, methods, SKU, Croston, method, seasonal, demand en forecast.

Allereerst wordt gekeken in de literatuur welke methodes er beschikbaar zijn en hoe deze tegen elkaar af te zetten zijn. Deze methodes worden ook beoordeeld op toepasbaarheid voor deze case.

1. Regression analyse ⁽¹⁷⁾
2. Exponential smoothing = a method of forecasting that uses smoothing factors to better estimate future demand for a SKU ^{(14) (15) (18) (19) (20) (21)}

- Has smoothing factors to compensate for variation factors and seasonal demand.
Bad with long inter-arrival time between orders.
3. Moving Average = a method of forecasting that guesses future demand on historical data. ⁽²¹⁾
Better under stable circumstances Has high smoothing of data since it only slowly changes the forecast. Forecast is slow to react in recognizing big changes in demand. Bad with long inter-arrival time between orders.
 4. Corston's Method takes account of both demand size and inter-arrival time between demands. The method is widely used in industry and it is incorporated in various best selling forecasting software packages ^{(13) (14) (15) (21) (22)}
 5. Syntetos–Boylan approximation ^{(14) (21) (22)}
 6. Neural Network = better under trending and seasonal variation, but not great with promotions ⁽¹⁴⁾
 7. Planning products that last only one season (fashion) ⁽¹⁶⁾

Bepaling aantal classes

Uit de literatuur volgt dat voor de bepaling van het aantal classes er 5 variabelen zijn die van belang zijn en hierop zal de case getoetst worden.

Variabelen

Het inzetten van een superfast zone is afhankelijk van enkele factoren. ⁽²³⁾

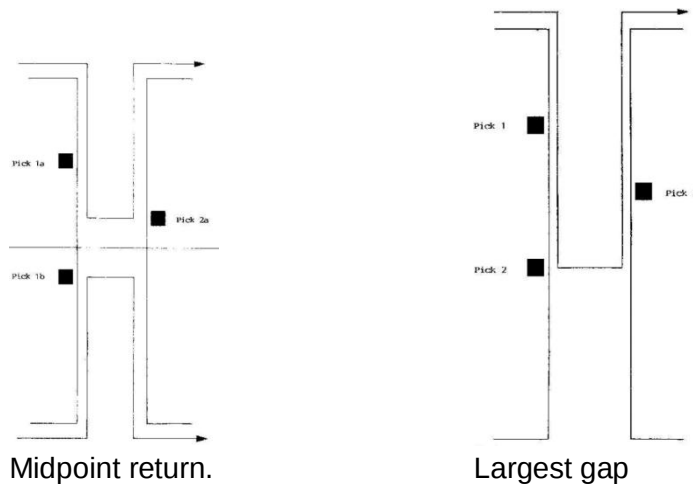
1. De som van het aantal order lines van de SKU superfast zone in verhouding tot het totaal aantal order lines en het fast aantal order lines.
2. Het aantal order lines per order.
3. Het verloop van superfast SKU. Als deze elke maand wisselen dan zal voorspelling lastig zijn en hierdoor is de zone minder nuttig worden.
4. Implementatiekosten vs opbrengsten.
5. Aantal picks per order line. Superfast vs fast vs slow

Vanuit deze 5 variabelen zal een conclusie getrokken worden of de toevoeging van een superfast zone toegevoegde waarde heeft voor de Wim Bosman Group en op grond hiervan zal een eindconclusie getrokken worden.

Berekening besparing

Om een goede inschatting te maken voor de gemaakte afstand van de route moet in acht genomen worden welke routing policy van toepassing is op deze case. Aan de hand van de theorie van Hall ⁽²⁴⁾ kan wordt er onderscheid gemaakt in 3 strategieën van routing.

- Traversal strategie, waarbij een aisle altijd volledig wordt overgestoken als een SKU zich in een aisle bevindt.
- Midpoint return strategie, waarbij de picker een pick doet vanaf de dichtbij zijnde kant en omkeert voor het midden.
- Largest gap strategie, waarbij de picker de picks doet zodat het de grootste rijtijd in een aisle vermeden wordt.



De case maakt duidelijk gebruik van het traversal algoritme bij de maken van de routes aangezien deze werkt met een S-shape routing en op grond hiervan is gekozen voor de berekening van hall die hierbij hoort.

Er is gekozen voor de Formule van Hall, omdat dit een formule is die relatief goede resultaten gegeven de beschikbare data. Tevens is deze formule robuust genoeg om gemiddelde aantallen mee te nemen.

Gezien de grote onzekerheid die gepaard is bij het voorspellen van de vraag bij gebruik van de planningstool is er besloten om dat de formule van Hall voldoende indicatie geeft voor de gebruikers om de resultaten van de planning te beoordelen.

De zwakte van de formule van Hall is dat deze niet geschikt is voor orders met heel weinig order lines in de situatie van deze case.

Formule van Hall mbt traversal algoritme

i = Zone (Fast of slow)

M_i = aantal (halve) aisles toegewezen aan zone i

N_i = gemiddeld aantal SKU per order toegewezen aan zone i

y = lengte (halve) aisle in meters

δ = gemiddelde gemaakte afstand per order

a = Aantal orders = totaal aantal order Lines / N_i

b = meter per sec gemiddeld afgelegd = snelheid wagen =

c = total time in uren

$$\delta = \sum_{i=F}^S M_i y \left(1 - \left(\frac{M_i - 1}{M_i} \right)^{N_i} \right)$$

$$c = \delta * a / (3600 * b)$$

Voorbeeld formule***Class based storage***

i = Zone = fast (F) of slow (S)

10 aisles

$M_F = 2$ aisles fast

$M_S = 8$ aisles slow

Een order heeft gemiddeld 10 order lines

80% wordt in slow gepickt

$$N_F = 10 * 0.8 = 8$$

20% wordt in fast gepickt

$$N_S = 10 * 0.2 = 2$$

Er zijn 1000 order lines in totaal

$$a = 1000/10 = 100$$

Een palletwagen rijdt 2 meter per seconde

$$b = 2$$

Lengte van een gang is 100 meter

$$y = 100$$

$$\delta = M_F y \left(1 - \left(\frac{M_F - 1}{M_F} \right)^{N_F} \right) + M_S y \left(1 - \left(\frac{M_S - 1}{M_S} \right)^{N_S} \right)$$

$$\delta = 2 * 100 \left(1 - \left(\frac{2-1}{2} \right)^8 \right) + 8 * 100 \left(1 - \left(\frac{8-1}{8} \right)^2 \right)$$

$$\delta = 199.2 + 187.5 = 386.7$$

$$c = \delta * a / (3600 * b)$$

$$c = 386.7 * 100 / (3600 * 2) = 5.3 \text{ uur}$$

Random storage

$$M_R = 10$$

$$N_R = 10$$

$$a = 100$$

$$b = 2$$

$$y = 100$$

$$\delta = 10 * 100 \left(1 - \left(\frac{10-1}{10} \right)^{10} \right) = 651.2$$

$$c = 9.0 \text{ uur}$$

Samentvatting literatuur

Storage assignment method	Literature
Turnover based, dedicated	(1) (2) (3) (4)
Random, closest open location	(1) (2) (3) (4) (12)
Class based	(1) (2) (3) (4)
Family based	(1)

Variabele voor toewijzing locaties	Literatuur
Family, similar SKU and similar size	(1) (4) (9)
Throughput	(1) (4) (7) (8) (9)
Demand frequency	(3) (4) (9) (10) (11)
Duration-of-stay	(4) (5) (8) (9)
Cube-per-order index	(2) (4) (6) (9)

Forecasting strategie	Location
Regression analyse	(17)
Exponential smoothing	(14) (15) (18) (19) (20) (21)
Moving Average	(21)
Corston's Method	(13) (14) (15) (21) (22)
Syntetos–Boylan approximation	(14) (21) (22)
Neural Network	(14)

Variabele voor aantal classes
Verhouding order lines
Aantal order lines per order
Verloop SKU
Aantal picks per order line
Implementatie kosten vs opbrengsten

De gekozen berekeningen voor de besparing is de formule van Hall die gebruik maakt van de Traversal strategie.

Hoofdstuk 4 : Keuzes

De literatuur wordt toegepast op deze case om zo een keuze te kunnen maken die geïmplementeerd wordt aan de hand van Excel.

Voor de elk van de gebieden waar onderzoek naar is verricht zal een keuze gemaakt worden voor deze case.

1. Storage assignment methode
2. Variabele voor toewijzing locaties
3. Forcasting strategie
4. Aantal classes
5. Berekening besparing

Storage assignment method

Storage assignment method	Savings	Gemak van implementatie	Onderhoudstijd en rekentijd
Turnover based, dedicated	+++	- -	- -
Random, closest open location	- - -	+++	+++
Class based	++	++	+
Family based	++	- -	-

Uit deze tabel blijkt dat **Class based storage assignment** op grond van savings, gemak van implementatie, onderhoud en rekentijd de beste optie is. Op alle criteria komt Class based storage goed naar voren.

Class based storage scoort het goed op Savings, omdat het relatief veel opbrengt 10/20% besparing is eenvoudiging haalbaar. De implementatie is relatief eenvoudig. Enkele simpele berekeningen geven al goede resultaten. De onderhoudstijd en rekentijd zijn goed in verhouding met de andere methodes waar berekening nodig is.

Random besparing komt ook goed naar voren aangezien het geen tijd en middelen kost, maar aangezien deze geen besparing oplevert is dit niet de moeite waard.

De overige opties leveren niet significant meer op en zijn veel duurder in implementatie en onderhoudt.

Variabele voor toewijzing locaties

Variabele voor toewijzing locaties	Combinatie met class based storage	Correlatie met gemaakte afstand	Implementatie en onderhoudskosten
Family, similar SKU and similar size	+/-	+++	-
Throughput	+++	+	+
Demand frequency	+++	+++	+
Duration-of-stay	- -	+++	- -
Cube-per-order index	+++	+/-	+/-

Op grond van deze resultaten is er gekozen om te werken met de variabele '**aantal order lines per maand**'. Het aantal order lines per maand = demand frequency.

Het aantal order lines is in directe correlatie met de gereden afstand, aangezien voor elke order line de route moet worden uitgebreid. De implementatie en onderhoudskosten zijn relatief beperkt en de demand frequency kan perfect gebruikt worden als variabele voor class bases storage.

Een bijkomend voordeel is dat het aantal order lines minder varieert bij acties en reclames. Dit komt voort uit het feit dat bij reclames en acties er vaak meer cases onder een order line vallen en dit houdt niet direct in dat er veel meer order lines zijn.

Forecasting strategie

Forecasting strategie	Savings	Implementatie/rekentijd/kosten	Toepasbaarheid Case
Regression analyse	+	+ +	++
Exponential smoothing	++	+	+/-
Moving Average	+	+ +	+/-
Corston's Method	+++	- -	-
Syntetos–Boylan approximation	++	- -	-
Neural Network	+++	- - -	- - -

Aangezien de meeste methodes niet geschikt zijn voor deze case, is er gekozen voor een variant van de **regression analyse met exponentiele gewichtstoekenning van de maanden**.

Er wordt gekeken naar data uit het verleden en hiervan wordt de beslissingsvariabele (order lines per maand) wordt bepaald. Vervolgens wordt met behulp van de gewichten van de maanden een gewogen gemiddelde bepaald en dit is de uiteindelijke beslissingsvariabele.

Aantal classes

Variabele voor aantal classes	Superfast/Fast/Slow	Fast/Slow
Verhouding order lines	- -	+
Aantal order lines per order	+/-	+/-
Verloop SKU	- -	++
Aantal picks per order line	-	+
Implementatie opbrengsten vs kosten	+/-	++

Hoofdstuk 5: Implementatie

In dit hoofdstuk wordt de implementatie van de case in Excel uiteengezet om stap voor stap een idee te geven hoe het uiteindelijke programma werkt. Tevens wordt de situatie na de implementatie kort beschreven.

Hierin zijn 3 onderdelen onderscheiden:

1. De data filtering en samenvoegingstool
2. De planningstool
3. Situatie na implementatie

Hoofdstuk 7: Conclusie

In opdracht van de Wim Bosman Group is er gekeken naar het verbeteren en eenvoudig updaten van de lay-out voor het warehouse. Om dit probleem aan te pakken is de volgende probleemstelling opgesteld:

Hoe kan de data uit het verleden worden gemodelleerd tot een advies advies over de classificatie van de SKU?

Om deze probleemstelling aan te pakken zijn de drie onderzoeksvragen aangepakt:

- *Welke variabelen zijn van belang voor het bepalen van een goede classificatie?*
Op grond van literatuuronderzoek en na overleg is gekozen voor de variabele:
Aantal order lines (per maand).
- *Hoe kan er rekening gehouden worden met de vervuiling van de data zoals reclames/ acties voor een SKU zodat geen verkeerde conclusies worden getrokken?*
Dit wordt ondervangen door te werken met periodes van een maand en het gebruik van de variabele 'aantal order lines'.
 - Het gebruik van maandelijkse periodes houdt in dat acties die een week duren minder opvallend in de data terugkomen.
 - Het 'aantal order lines' is een variabele die minder invloed ondervindt van acties en reclames door het feit dat het aantal casepicks hierin niet terugkomt. Het aantal casepicks zal omhoog gaan met een reclame (meer dozen verkocht), maar de order lines zullen hier geen directe invloed door hebben aangezien er voldoende cases op een pallet liggen. (over het algemeen gaan er 5 tot 100 order lines uit een pallet). Zelfs een order met waarbij 1 SKU heel veel cases bevat zal bijna nooit van meer dan 2 pallets moeten picken aangezien deze anders uit de bulk gehaald worden.

Door deze twee keuzes zal de variantie in de vraag weinig reactie tonen op reclames en acties.

Naast de vervuiling van acties en reclame was er nog een meer directe vervuiling in de vorm van SKU's, die wel in de data voorkomen, maar geen onderdeel worden van een van de 4 zones. Om deze vervuiling tegen te gaan wordt is het mogelijk deze uit de data te filteren voordat een planning gemaakt wordt. Dit wordt met behulp van macro's in Excel, automatisch gedaan.

- *Welke methodes van terug/vooruit kijken zijn geschikt voor de indeling in slow/fast?*
Deze onderzoeksvraag was lastig direct uit de literatuur te halen, omdat er veel verschil in de gebruikte methodieken is. De twee overheersende methodieken zijn corston's method en exponential smoothing. Beide methodieken worden veel gebruikt in het bedrijfsleven en worden vaak geleverd als planning tools in een ERP programma. Er is veel data voor nodig om een goede inschatting te krijgen. In deze case variëren de SKU's jaarlijks sterk en de SKU's zijn ook seizoensgebonden. Dit maakt het lastig ze in te schatten met deze twee methodieken.
Tevens is de meeste literatuur met betrekking tot voorspellingen niet gericht op het

classificeren van SKU, maar de literatuur is gericht op voorspellen van exacte voorraad.

Vanuit deze optiek is gekozen om een variant op de regression analyse te gebruiken voor de berekening van een gewogen gemiddelde. Er wordt een dataset (aantal maanden) bepaald als input. De maanden zullen een gewicht toegedeeld krijgen op grond van het feit hoe ver de maanden van de target maand af liggen.

Voorbeeld: target maand is mei met 3 maanden data:

Resultaat: februari 25%, maart 50% en april 100%

De gebruiker kan in de tool er ook voor kiezen zelf in te stellen welke gewichten alle maanden krijgen en daarmee de resultaten te beïnvloeden.

Op grond van de uitkomsten van de onderzoeksvragen is een tool geprogrammeerd in Excel die het mogelijk maakt snel en eenvoudig een advies te genereren voor de Classificatie van alle SKU's. Deze tool gebruikt de data die uit het warehouse management systeem en de input van gebruiker over de parameters en variabelen (target maand, maanden terug/vooruit kijken, beslissingsvariabelen) om elke SKU aan een categorie toe te wijzen.

Uit initiële testen blijkt dat de tool een betrouwbaarheid heeft van meer dan 90% heeft met betrekking tot voorspelling van fast movers als de default waardes worden aangehouden. Dit houdt in een besparing van ongeveer 35 uur ten opzichte van een random indeling.

Aanbevelingen

Op grond van dit project zijn er nog enkele aanbevelingen en adviezen voor vervolgprojecten.

1. Het is optimaal de planning zo vroeg mogelijk in de maand te draaien. Dit omdat de data recentelijk is geüpdate (Orders delivered) en hierdoor zal de impact van de herindeling de grootste invloed hebben en het langste nut (een volle maand). Aangezien de CpC en SkuDim wel op de dag van het plannen worden gemaakt zullen hierin inconsistenties staan tegenover de OD. Zoals nieuwe SKU en voorraad die niet in verhouding is met de verwachte vraag (die 2^e is alleen van invloed op old en new SKU).
2. De besparing die voorkomt uit het maken van een fast/slow indeling zal wegvallen als er niet elke maand een update gehouden wordt om de inrichting bij te werken. Veel van de SKU's variëren sterk in vraag in de loop van het jaar en zullen dus regelmatig een andere categorie krijgen. Als er te lang geen update is geweest kan een random indeling zelfs beter zijn dan een class indeling. Als er heel veel fast movers in de slow zone liggen zal het meer tijd kosten deze te bereiken dan bij een random indeling. Dit kan het geval zijn bij seizoensgebonden SKU's als de planning in de zomer is gemaakt en het nu winter is.
3. Om de rijtijdsbesparing nog verder te verbeteren kan er gekeken worden naar de verwantschappen tussen de orders. Als er veel orders zijn die dezelfde SKU's bevatten, kan het nuttig zijn om deze SKU's naast elkaar in het warehouse te leggen om zo een tijdsbesparing te bereiken. Dit is een zeer rekenintensief project met mogelijkheden voor grote besparingen met behulp de juiste heuristieken.
4. Naast dit project is er ook nog een verbeteringslag mogelijk om de tijdsduur bij multiple locations te verminderen. Het komt relatief vaak voor dat een pallet leeg raakt en dit houdt in dat er meestal een order picker een multiple location voor deze

SKU heeft. Dit houdt in dat de order picker stopt met de normale route en daarna heen en weer moet rijden naar de 2^e locatie van de SKU. Om de tijd hiervan te beperken is het handig om de 2^e locatie zo dicht mogelijk bij de huidige locatie te plaatsen.

Door het huidige project zal een fast SKU altijd in de fast zone blijven liggen en hiermee wordt een SKU in ieder geval relatief dicht bij de huidige locatie geplaatst. Maar toch kan het gunstig zijn om te kijken of het 'move algoritme' voor het plaatsen van nieuwe pallets geoptimaliseerd kan worden.

Het doel van het move algoritme zou moeten zijn een pallet op de dichtstbijzijnde open plek te plaatsen. De besparing die hiermee bereikt kan worden is waarschijnlijk groter dan een class indeling kan bereiken, aangezien het omrijden binnen de route veel tijd kost elke keer.

5. Op grond van de tijd op de werkvloer viel mij op dat er tijd verloren gaat aan het omstapelen van pallets, omdat zware SKU op lichte SKU staan.
Dit komt voort uit het feit dat Heavy vs Light op dit moment de kritieke grens van 3,5 kilo heeft. Maar kleine dozen die net geen 3,5 kilo wegen worden als Light geclassificeerd. Maar deze dozen wegen relatief meer dan grote dozen die net 3,5 kilo wegen en hierdoor moeten soms dozen omgestapeld worden.
Vandaar dat mijn advies is om het Heavy/Light criterium niet vast om 3,5 te leggen maar deze te laten variëren met het oppervlakte (niet inhoud) van de dozen en eventueel het materiaal van het karton van de doos.
 - Dozen met een lichter gewicht zullen eerder als light geclassificeerd worden.
 - Dozen met groter oppervlak zullen eerder als light geclassificeerd worden.
 - Dozen met dunner karton zullen eerder als light geclassificeerd worden, aangezien ze door kunnen zakken.

Hoofdstuk 8: Gebruikte literatuur

- ⁽¹⁾ Batch Construction Heuristics and Storage Assignment Strategies for Walk/Ride, and PickSystems Robert A. Ruben • F. Robert Jacobs (1999)
- ⁽²⁾ A branch and bound algorithm for class based storage location assignment Venkata Reddy Muppani (Muppant), Gajendra Kumar Adil* (2007)
- ⁽³⁾ A comparison of picking, storage, and routing policies in manual order picking Charles G. Petersen, Gerald Aase (2004)
- ⁽⁴⁾ Research on warehouse operation: A comprehensive review Jinxiang Gu, Marc Goetschalckx, Leon F. McGinnis. (2007) Zie table 4 voor meer artikelen mbt dynamic storage
- ⁽⁵⁾ A review of warehouse models, Gilles Cormier, Eldon A. Gunn (1992)
- ⁽⁶⁾ A study of storage assignment problem for an order picking line in a pick-and-pass warehousing system, Jason Chao-Hsien Pan, Ming-Hung Wu (2009)
- ⁽⁷⁾ New batch construction heuristics to optimise the performance of order picking systems, Ling-Feng Hsieh, Yi-Chen Huang (2011)
- ⁽⁸⁾ Research on warehouse design and performance evaluation: A comprehensive review, Jinxiang Gu, Marc Goetschalckx, Leon F. McGinnis (2010)
- ⁽⁹⁾ Storage location assignment: Using the product structure to reduce order picking times, H. Brynz, M.I. Johansson (1996)
- ⁽¹⁰⁾ Storage location assignment and order picking optimization in the automotive industry Seval Ene, Nursel Öztürk (2010)
- ⁽¹¹⁾ Warehouse design optimization, J. Ashayri and L.F. Gelders (1985)
- ⁽¹²⁾ A study of storage assignment problem for an order picking line in a pick-and-pass warehousing system, Jason Chao-Hsien Pan, Ming-Hung Wu (2009)
- ⁽¹³⁾ On the bias of Croston's forecasting method. Ruud Teunter, Babangida Sani (2009)
- ⁽¹⁴⁾ Lumpy demand forecasting using neural networks, Rafael S. Gutierrez, Adriano O. Solis, Somnath Mukhopadhyay (2007)
- ⁽¹⁵⁾ On the empirical performance of (T,s,S) heuristics, M. Zied Babaia, Aris A. Syntetos, Ruud Teunter (2009)
- ⁽¹⁶⁾ Forecasting demand for single-period products: A case study in the apparel industry, Julien Mostard, Ruud Teunter, René de Koster (2010)
- ⁽¹⁷⁾ SKU demand forecasting in the presence of promotions, Özden Gür Ali, Serpil Sayin, Tom van Woensel, Jan Fransoo (2009)
- ⁽¹⁸⁾ Exponential smoothing: The state of the art—Part II, Everette S. Gardner Jr.

- ⁽¹⁹⁾ Forecasting for inventory control with exponential smoothing, Ralph D. Snyder , Anne B. Koehler , J. Keith Ord
- ⁽²⁰⁾ Exponential smoothing models: Means and variances for lead-time demand, Ralph D. Snyder, Anne B. Koehler, Rob J. Hyndman, J. Keith Ord
- ⁽²¹⁾ The accuracy of intermittent demand estimates, Aris A. Syntetos, John E. Boylan
- ⁽²²⁾ Forecasting-based SKU classification, G. Heinecke, A.A. Syntetos, W. Wang
- ⁽²³⁾ Improving order-picking performance through the implementation of class-based storage, Charles G. Petersen, Gerald R. Aase, Daniel R. Heiser (2004)
- ⁽²⁴⁾ Distance approximations for routing manual pickers in a warehouse R.W.Hall (1993)
- ⁽²⁵⁾ Lecture sheets Warehousing University Twente (2012)