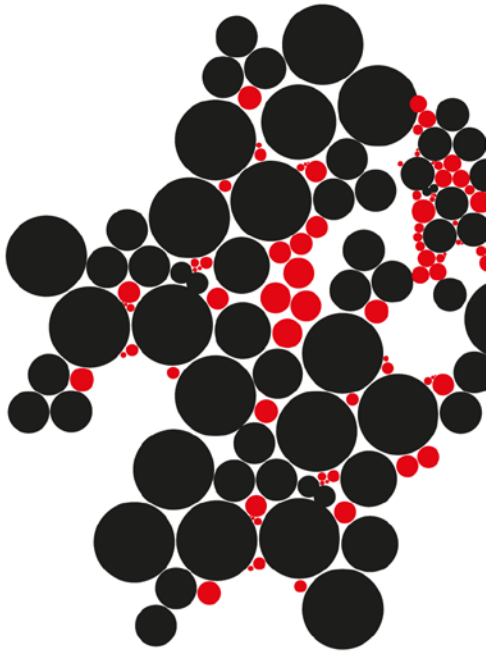
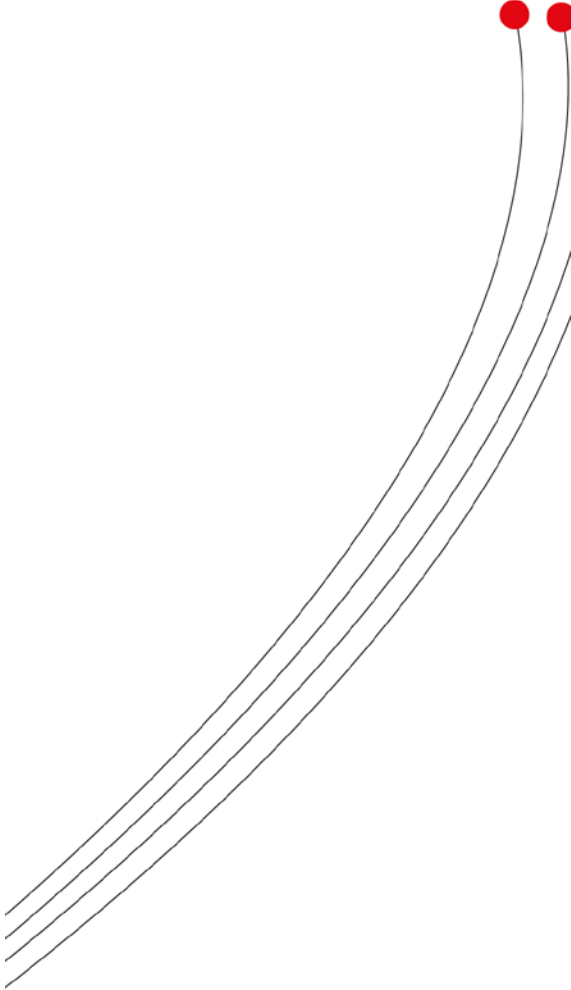




DEMAND FORECASTING FOR SPARE PARTS AT STORK



Berry Gerrits

TECHNISCHE BEDRIJFSKUNDE
UNIVERSITEIT TWENTE

BEGELEIDER
Matthieu van der Heijden

STATUS
EINDVERSLAG

Bachelorscriptie

Technische Bedrijfskunde



Food & Dairy Systems

<u>Externe begeleider</u> Massimo Schutte Stork Food & Dairy Systems B.V. Deccaweg 32, 1042 AD AMSTERDAM P.O. Box 759, 1000 AT AMSTERDAM The Netherlands Tel.: +31 (0) 20 6348 506 E-mail: massimo.schutte@sfds.eu	<u>Interne begeleider</u> Matthieu van der Heijden Universiteit Twente Drienerlolaan 5, 7500 AE ENSCHEDE Ravelijn 3357 The Netherlands Tel.: +31 (53) 48 928 52 E-mail: m.c.vanderheijden@utwente.nl
<u>Extern contactpersoon</u> Maarten Driessen Gordian Logistic Experts B.V. Groenewoudsedijk 63, 3528 BG UTRECHT The Netherlands Tel.: +31 (0) 30 6866 980 E-mail: m.driessen@gordian.nl	<u>Student</u> Berry Gerrits Universiteit Twente Calslaan 60-91, 7522 MG ENSCHEDE The Netherlands Tel.: +31 (6) 13 3463 55 E-mail: s1005367@utwente.nl

Voorwoord

Dit verslag is het resultaat van een drie maand durende afstudeerstage bij Stork Food & Dairy Systems (SFDS) te Amsterdam voor mijn bacheloropdracht van Technische Bedrijfskunde aan de Universiteit Twente. Deze opdracht is gericht op het inzicht verwerven van vraagkarakteristieken van de spare parts die Stork levert aan haar klanten. Aan de hand van deze karakteristieken moeten geschikte voorspelmethodeken gekozen worden om zo tot een juiste en betrouwbare voorspelling voor toekomstige periode(n) te komen. Deze informatie is belangrijke input voor de voorraadbeheer bij Stork. Aan de hand van de voorspelling (en andere input) kan de voorraad goed gemanaged worden en hierdoor geen onnodig hoge voorraad te hebben of juist vaak te maken te hebben met stock-outs. Het uiteindelijke doel van Stork is om met zo weinig mogelijk voorraad een zo hoog mogelijke beschikbaar van spare parts te generen.

Dit onderzoek was nooit in deze hoedanigheid gerealiseerd zonder de belangrijke input van vele mensen. Daarom wil ik allereerst mijn interne begeleider, Matthieu van der Heijden, bedanken voor het voorleggen van deze opdracht. Daarnaast was zijn feedback van zeer toegevoegde waarde en hij heeft dit verslag naar een hoger niveau getild. Ook wil ik Maarten Driessen van Gordian bedanken voor ons gesprek en zijn aanbeveling van mij bij Stork. Daarnaast wil ik mijn externe begeleider bij Stork, Massimo Schutte, bedanken voor de begeleiding die hij mij heeft gegeven de afgelopen maanden. De samenwerking verliep uitstekend en het bedrijf Stork verdient een compliment voor het begeleiden van een afstudeerder. Ook wil ik Anne Dirk Goos en Hans Gaillard van Stork vriendelijk bedanken voor de waardevolle discussies die we met elkaar hebben gehad. Het verschaffen van de vele te analyseren data heb ik te danken aan Ron van Oostrum, hij heeft mij uitstekend geholpen bij het verkrijgen van data, waarvoor mijn dank. Ook de vele tegenlezers die dit verslag onder ogen hebben gehad, wil ik graag bedanken voor hun tijd en moeite.

Op de laatste pagina van dit verslag is een leeswijzer toegevoegd. Deze kan worden uitgekapt en zo naast het verslag gebruikt worden. Het gebruik van de leeswijzer is niet noodzakelijk, maar voor de lezer met minder voorkennis van het onderwerp kan het van pas komen.

Management samenvatting

Stork Food & Dairy Systems (SFDS) te Amsterdam ontwikkelt, produceert en levert geïntegreerde systemen voor het verwerken en afvullen van o.a. zuivel en sappen. Deze kapitaalintensieve machines maken deel uit van het primaire productieproces van de klant en dienen dus een zo hoog mogelijke up-time te hebben. Om deze up-time te garanderen levert SFDS o.a. spare parts aan haar klanten. Het inzicht in de vraag naar spare parts en het houden van de juiste voorraad voor een optimaale beschikbaarheid is momenteel nog onvoldoende. Daarom richt dit onderzoek zich op de vraagvoorspelling op artikelniveau van de spare parts die SFDS aan haar klanten levert. Allereerst is er gekeken naar de huidige manier van vraagvoorspelling en diens kwaliteit. Daarna is onderzoek gedaan om de spare parts te classificeren voor het kiezen van een juiste voorspelmethode en vervolgens zijn er methodieken gezocht en getest.

De huidige manier van vraagvoorspelling bestaat uit de gemiddelde vraag van de afgelopen vier jaar. Met dit gegeven per artikel, worden de parameters bepaald voor het voorraadbeheer. Dit geeft echter geen inzicht in de karakteristieken van de vraag. Om deze karakteristieken beter in beeld te brengen, is een theoretisch kader opgesteld die de spare parts op basis van twee variabelen classificeert: 1) ADI, de gemiddelde tussentijd tussen vraagmomenten en 2) CV^2 , de gekwadrateerde variatiecoëfficiënt; de gekwadrateerde verhouding tussen de standaardafwijking en het gemiddelde van de vraaghoeveelheid. Als beide variabelen laag zijn, dan wordt dit een 'smooth' artikel genoemd. Bij een hoge CV^2 en een lage ADI, dan heet dit een 'erratic' artikel. Bij een hoge ADI, maar een lage CV^2 , dan heet dit een 'intermittent' artikel. En als allebei de variabelen hoog zijn, dan heet dit een 'lumpy' artikel. Op basis van deze classificatie en literatuuronderzoek is gekozen om in de 'smooth' categorie de methodiek van Croston toe te passen en in de andere categorieën de Syntetos & Boylan Approximation.

De prestaties van de methodieken zijn gemeten aan de hand van verschillende prestatie-indicatoren (Mean Squared Error, Adjusted-Mean Absolute Percentage Error (A-MAPE), Cumulated Forecast Error, Number of Shortages en Periods in Stock) en hieruit bleek dat de smooth categorie het best te voorspellen is (met een gemiddelde A-MAPE van 55,93%), vervolgens erratic (71,93%), daarna intermittent (103,80%) en tenslotte lumpy (127,18%). Er is gekozen voor de A-MAPE in plaats van de MAPE omdat de eerstgenoemde om kan gaan met nul-vraag door de gemiddelde afwijking van de voorspelling te delen door de gemiddelde daadwerkelijke vraag.

De prestaties van de methodieken waren echter niet significant beter (in termen van MSE zelfs iets slechter, bij A-MAPE en PIS waren er geen significante verschillen) dan de huidige manier van voorspellen. De voorspelmethode zijn echter wel dynamischer dan een vierjaargemiddelde en de verwachting is dan ook dat de voorspelmethodeken betere input geven voor het voorraadbeheer. Het voorraadbeheer op basis van de voorspelling dient echter nog onderzocht te worden.

Behalve de analyse op basis van historische vraaggegevens is ook gekeken naar andere informatiebronnen. Het bleek dat de voorspelaccuratie enorm verbetert als er andere (deterministische) informatie in de voorspelling gebracht door bijvoorbeeld het gebruik van (deels) van te voren bekende vraag. Dit kan op basis van bekende onderhoudsplannen en overleg met de klant bewerkstelligd worden.

Om het proces van vraagvoorspelling voor spare parts dynamisch, snel en effectief te gebruiken binnen Stork wordt aanbevolen om een softwarepakket als Slim4 te gebruiken.

Inhoudsopgave

Voorwoord	IV
Management samenvatting	V
Hoofdstuk 1 Inleiding Stork	1
Hoofdstuk 2 Inleiding onderzoek	2
Aanleiding onderzoek	2
Inleiding spare parts	2
Probleemstelling	3
Onderzoeksdoel	3
Deelvragen	3
Hoofdstuk 3 Huidige situatie	4
Huidige manier van vraagvoorspelling	4
Toetsing van huidige vraagvoorspelling	5
Prestatie-indicatoren	5
Resultaten huidige voorspelling	8
Hoofdstuk 4 Spare parts classificatie	9
Forecast Support System	9
Theoretisch kader	9
Scheidingswaarden	11
Hoofdstuk 5 Voorspelmethodeken	12
Literatuuroverzicht	12
Geschikte methodeken	12
Gekozen methodeken	15
Croston's method (CR)	15
Syntetos & Boylan Approximation (SBA)	16
Moving Average (MA)	16
Hoofdstuk 6 Data analyse	17
Aanpak data analyse	17
Resultaten data analyse	17
Homogeniteitsanalyse	18
Aanpak vervolgvraag	20
Hoofdstuk 7 Empirische test	22
Prestatie-indicatoren	22
Test-set	22
Aanpak	22

Resultaten	22
Categorie: Smooth, Voorspelmethodiek: Croston.....	23
Categorie: Erratic, voorspelmethodiek: SBA.....	25
Categorie: Intermittent, voorspelmethodiek: SBA en MA.....	26
Categorie: Lumpy, voorspelmethodiek: SBA en MA.....	27
Vergelijking huidige voorspelling met voorspelmethodiek	30
Hoofdstuk 8 Het aanpassen van de voorspelling	31
Aanpassingen aan de voorspelling.....	31
Geïntensifieerd klantcontact.....	31
Orders inleggen	32
Voorbeeld orders inleggen.....	32
Hoofdstuk 9 Consequenties voor voorraadbeheer	34
Interactie voorspelling en voorraadbeheer	34
Invloed op voorraadbeheer	35
Hoofdstuk 10 De relatie tussen klantinformatie en voorspelbaarheid.....	36
Installed base forecasting	36
Hoofdstuk 11 Conclusies en aanbevelingen.....	37
Aanbevelingen.....	38
Bibliografie	39
Bijlagen	42
Bijlage 1 – Prestatie-indicatoren	43
Bijlage 2 – Overzicht voorspelmethodiek	45
Bijlage 3 – Uitwerking bias-gecorrigeerde factor SBA.....	48
Bijlage 4 – Data analyse (1)	49
Bijlage 5 – Data analyse (2)	50
Bijlage 6 – Uitwerking Syntetos-Boylan Approximation	51
Bijlage 7 – Grafische weergave per categorie.....	53
Categorie: Smooth, Voorspelmethodiek: Croston.....	53
Categorie: Erratic, voorspelmethodiek: SBA.....	55
Categorie: Intermittent, voorspelmethodiek: SBA	57
Categorie: Lumpy, voorspelmethodiek: SBA	59
Bijlage 8 – Overzicht resultaten van alle artikelen.....	61
Bijlage 9 – Voorbeelden invloed op voorraadbeheer	66
Bijlage 10 – Lijst met figuren en tabellen.....	67

Hoofdstuk 1 Inleiding Stork

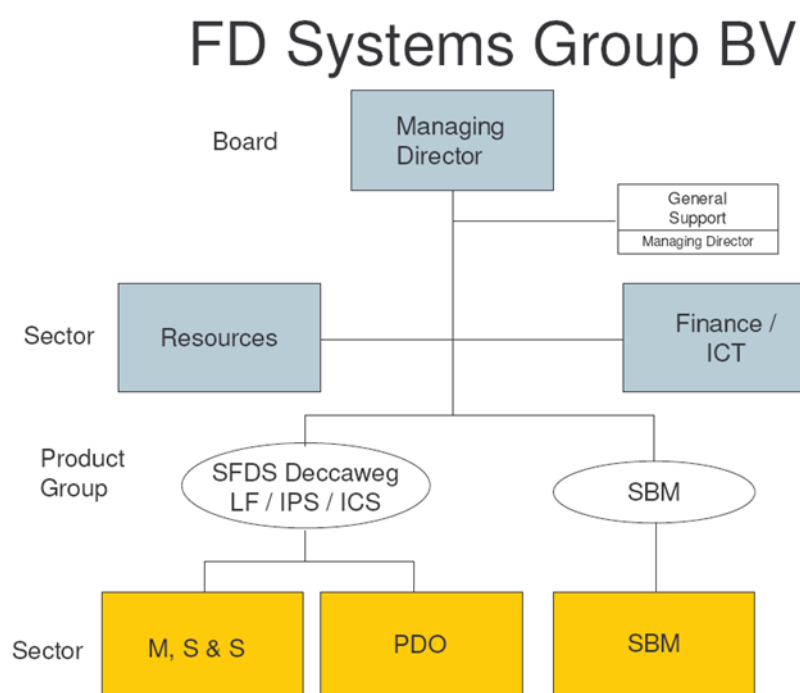
Dit hoofdstuk behandelt kort wat algemene informatie over Stork Food & Dairy Systems.

Stork Food & Dairy Systems (SFDS) ontwikkelt, produceert en levert geïntegreerde systemen voor de verwerking en afvulling van o.a. zuivel en sappen. Deze kapitaalintensieve systemen worden wereldwijd geleverd en after-sales service wordt dan ook in heel de wereld verleend. Het hoofdkantoor van SFDS staat in Amsterdam en daar worden lineaire aseptische vulsystemen (LF), in-proces sterilisatie (IPS) systemen en in-container sterilisatie (ICS) systemen ontwikkeld, geproduceerd en geleverd aan klanten. Deze systemen worden geleverd bij (grote) klanten waar ze deel uitmaken van het primaire productieproces. Vanwege de kritiekheid van de systemen voor het productieproces van de klanten wil SFDS een zo hoog mogelijke up-time van haar systemen leveren.

Stork is van origine een Nederlands bedrijf en levert sinds de jaren '30 van de vorige eeuw systemen voor de food & dairy markt. In 2006-2007 werd het beursgenoteerde Stork opgesplitst en ging verder onder leiding van verschillende investeringsfondsen. Voor de opsplitsing bestond Stork uit de volgende vier divisies:

- Textieldrukmachines (Stork Prints)
- PMT (Kippenslachtmachines), Townsend (roodvleesverwerking) en Food & Dairy systems (Stork Food Systems)
- Lucht- en ruimtevaartdivisie (Fokker Aerospace Group)
- Technische dienstverlening (Stork Industry Services)

In 2010 nam de Nederlandse investeringsmaatschappij Nimbus SFDS inclusief SBM (Stork Bottling Maintenance) en ICS (In Container Sterilisatie) over. SFDS richt zich met name op de markt van het industrieel lang houdbaar maken van voedingsmiddelen in consumentenverpakkingen. Met name de verwerking van voeding, zuivel en vruchtensappen zijn marktsegmenten die belangrijk zijn voor Stork. De organisatiestructuur van de Food Systems Group staat hieronder afgebeeld.



Hoofdstuk 2 Inleiding onderzoek

Dit hoofdstuk gaat in op de inhoud van het onderzoek. Allereerst wordt de aanleiding beschreven, gevolgd door een korte introductie over spare parts en tenslotte de probleemstelling, het onderzoeksdoel en enkele deelvragen.

Aanleiding onderzoek

Om een hoge up-time te garanderen naar de klanten en om bepaalde service levels te borgen is het ook belangrijk om het proces van reserve onderdelen (spare parts) goed te beheersen. Uitgangspunt is de beschikbaarheid van de juiste onderdelen voor zowel preventief als correctief onderhoud op het juiste moment. Daarbij zijn spare parts een belangrijke en groeiende onderstroom van de omzet van SFDS. Ook als er een breakdown komt bij een klant (vb. een kritiek onderdeel van het systeem gaat kapot) is het noodzaak deze tijdig (lees: zo snel mogelijk) te vervangen gezien de hoge kosten die gemoeid zijn met het stilvallen van het productieproces bij de klant. Gezien de lange levertijd van (klantspecifieke) onderdelen is het essentieel om de juiste spare part op het juiste moment beschikbaar te hebben, rekening houdend met de kosten die gemoeid zijn met het houden van voorraad. Daarnaast is het belangrijk dat een klant op tijd preventief onderhoud kan plegen en dat de daarbij behorende onderdelen ook beschikbaar zijn. Mochten één of meerdere artikelen niet beschikbaar zijn dan wordt het onderhoud uitgesteld en dit is slecht voor het imago van Stork en de klanttevredenheid. Hierdoor bestaat de kans dat een klant niet langer onderdelen bij Stork bestelt, maar via andere wegen. Dit brengt ook nog inkomstenderving met zich mee.

Om dit proces te beheersen is het allereerst noodzakelijk om een goede grip op de vraag naar spare parts te krijgen. Momenteel heeft Stork moeite met het beheersen van de vraag naar spare parts.

Inleiding spare parts

In tegenstelling tot bijvoorbeeld onderdelen van nieuwbouw is het managen van de levering van spare parts vaak lastig en complex. Dit komt met name door de aard van het product. Een spare part wordt namelijk gebruikt voor het repareren of het onderhouden van een machine. Doordat het vaak niet in te schatten is wanneer een onderdeel stuk gaat, is het ook lastig om de vraag hiernaar in te schatten. Daarnaast kan het voorkomen dat ten behoeve van groot onderhoud of voorraadaanvulling van de klant ineens grote hoeveelheden spare parts worden besteld. Dus de vraag naar spare parts wordt over het algemeen gekarakteriseerd door een sporadische vraag (veel perioden hebben geen vraag) en als er dan vraag is, kan deze zeer gevarieerd zijn (Boylan & Syntetos, 2010). Door deze karakteristieken is het moeilijk om op een kostenefficiënte wijze het juiste onderdeel op het juiste tijdstip beschikbaar te hebben. Daarnaast zijn sommige spare parts erg duur in aanschaf en/of hebben een lange levertijd. Door de hoge kosten die gemoeid zijn bij het stilvallen van een machine bij de klant is het essentieel dat de goede spare parts snel geleverd kunnen worden. Niet alleen vanuit de after sales service voor de klant zijn spare parts belangrijk, maar over het voor Stork zijn spare parts ook erg lucratief. Spare parts management is dus erg belangrijk en raakt diverse afdelingen binnen het bedrijf.

Alvorens dit onderzoek gestart is, is er het nodige werk gedaan binnen SFDS naar het assortiment aan spare parts. De Product Manager heeft zo'n 6000 SKU's (Stock Keeping Units) geclassificeerd als actief onderhouden service artikel qua prijs, (technische) informatie en logistiek. Een actief service artikel wil zeggen dat het de afgelopen drie jaar minstens twee keer verkocht is. Er wordt onderscheid gemaakt tussen kritische onderdelen (critical) en onderhoudsonderdelen (maintenance). Voortbordurend op dit interne onderzoek wordt hier gefocust op het aspect van

vraagvoorspelling. De scope van dit onderzoek is dan vooral de vraagkant en niet het voorraadbeheer wat hieruit zou moeten voortvloeien. De uitkomsten van dit onderzoek zullen input zijn voor een onderzoek dat zich richt op voorraadbeheer.

Probleemstelling

Er is onvoldoende inzicht in de vraagpatronen van spare parts en daardoor is er geen optimale voorraad en beschikbaarheid.

Onderzoeksdoel

Het zo goed mogelijk in kaart brengen van de vraag naar diverse categorieën spare parts door de toepassing van voorspelmethodeken zodat er per artikel een maandelijkse prognose kan worden bepaald.

Deelvragen

Om het onderzoeksdoel te behalen zijn de volgende deelvragen gedefinieerd

1. Hoe is de huidige wijze van vraagvoorspelling en wat is de kwaliteit hier van?
2. Hoe kunnen spare parts op een slimme manier worden gecategoriseerd ten behoeve van vraagvoorspelling?
3. Welke voorspellingsmethodeken zijn geschikt voor de diverse categorieën spare parts?
4. Hoe kunnen geselecteerde voorspellingsmethodeken getoetst worden?
5. Hoe kunnen de geselecteerde voorspellingsmethodeken worden onderhouden?

Om dit onderzoek verder vorm te geven, wordt in hoofdstuk 3 beschreven hoe de huidige manier van voorspellen is en hoe deze presteert. Vervolgens wordt in hoofdstuk 4 een classificatie voorgesteld om de spare parts te categoriseren ten behoeve van vraagvoorspelling en wordt er in hoofdstuk 5 gekeken naar voorspelmethodeken voor deze categorieën. In hoofdstuk 6 wordt een data-analyse uitgevoerd en daarna wordt in hoofdstuk 7 gekeken naar de prestaties van de gekozen methodek(en). Vervolgens komt er een hoofdstuk over mogelijke aanpassingen aan de voorspelling.

Vanwege extra tijd is dit onderzoek uiteindelijk iets uitgebreid. Daarom is in hoofdstuk 9 een voorzichtige blik geworpen naar voorraadbeheer en is in hoofdstuk 10 gekeken naar hoe verbeterd klantcontact invloed kan hebben op de voorspelbaarheid. Hoofdstuk 11 geeft de conclusies en aanbevelingen.

Hoofdstuk 3 Huidige situatie

Dit hoofdstuk beschrijft de huidige manier van vraagvoorspelling. Daarna wordt de huidige manier van voorspellen getoetst aan de hand van enkele prestatie-indicatoren.

Huidige manier van vraagvoorspelling

Momenteel wordt er jaarlijks gekeken hoeveel omzet er gerealiseerd kan worden door het verkopen van spare parts. Deze omzet wordt berekend aan de hand van de installed base en het marktaandeel. De te verwachten vraag wordt dan gesplitst in twee categorieën: historisch en onderhoudsplan. In onderhoudsplannen is zeer gedetailleerd uitgewerkt welk onderhoud een bepaalde machine nodig zou moeten hebben om zo de machine draaiende en betrouwbaar te houden. Deze onderhoudsplannen zijn gedefinieerd voor machines van na 2005. In deze plannen zijn lijsten met spare parts gedefinieerd die gebruikt moeten worden tijdens deze maintenancebeurten. Op basis van de installed base wordt gekeken hoeveel van deze maintenance beurten het komende jaar voor kunnen gaan komen. Op basis van dit gegeven kun je een totaal lijst per jaar krijgen met benodigde spares en eveneens de te verwachten omzet die daar uit gegenereerd kan worden. Dit is echter een ideale wereld scenario aangezien niet alle klanten ook daadwerkelijk deze voorgeschreven plannen volgen of (deels) spare parts bij een andere leverancier bestellen. Daarom wordt door de Area Service Manager ingecalculeerd hoeveel marktaandeel Stork heeft binnen het gebied waar zij of hij verantwoordelijk voor is. Op basis hiervan wordt een percentage gecalculeerd dat aangeeft hoeveel klanten daadwerkelijk deze spare parts gaan bestellen. Zo kan worden bepaald hoeveel vraag per artikel per jaar verwacht kan worden. Dit is momenteel hoe de vraag naar spare parts voorspeld wordt. Het nadeel van deze methode is dat het nogal grof is. De voorspelling per jaar blijkt geen goede graadmeter voor de vraag naar spares gedurende het jaar omdat uit de praktijk blijkt dat deze niet uniform over het jaar verdeeld is. Daarom kan bijvoorbeeld gekozen worden om de complete ingeschatte behoefte aan het begin van het jaar op voorraad te leggen, maar dit brengt veel kosten met zich mee. Aan de andere kant is een uniforme benadering ook geen representatieve weergave van de werkelijkheid.

Als een onderdeel niet gedefinieerd is in een onderhoudsplan, dan wordt deze geclassificeerd als historisch. Dit wil zeggen dat het artikel alleen voorkomt in machines ouder dan 2005. Wel moet opgemerkt worden dat onderhoudsplannen natuurlijk ook op ervaring en historie zijn gebaseerd. Als een onderdeel dus niet in een onderhoudsplan zit (oudere onderdelen) worden deze dus behandeld als historisch. Hier wordt simpelweg gekeken naar het gemiddelde verbruik van de afgelopen vier jaar.

Zoals te zien is de huidige aanpak nogal grof (per jaar), terwijl de karakteristieken van spare parts juist op een kleiner interval beter te zien zijn. Daarom is een analyse van historische gegevens van toevoegde waarde op de huidige methodiek. Desondanks is deze manier van aanpak wel erg klantgericht en houdt dus ook rekening met bewegingen vanuit de markt (bijvoorbeeld marktaandeel), wat een historische analyse niet doet. Daarom is het uiterst zinvol om deze twee methodieken uiteindelijk bij elkaar te leggen om zo tot een goed oordeel te komen over wat te verwachten en wat realistisch is.

Toch is de mening binnen het bedrijf dat de voorspelling op jaarbasis, qua aantal, redelijk goed is. Het verschil tussen de verwachte afzet per jaar en de daadwerkelijke vraag zou dus aardig laag moeten

zijn. Toch is de vraag gedurende het jaar slecht in kaart gebracht. De analyse van de vraagpatronen moet hier invulling aan geven.

Toetsing van huidige vraagvoorspelling

Zoals gezegd wordt er momenteel alleen per jaar voorspeld. Daarnaast wordt er van een uniforme verdeling uitgegaan. Om aan het eind van het onderzoek de huidige situatie met de nieuwe te vergelijken, zijn de jaarvoorspellingen omgezet in maandvoorspellingen door de jaarvoorspelling te delen door 12 (uniforme verdeling). De dataset (62 artikelen) die gebruikt wordt, is dezelfde dataset als in hoofdstuk 7 gebruikt gaat worden. De voorspelling van 2011 zal getoetst worden aan de hand van enkele prestatie-indicatoren zoals beschreven in de volgende paragraaf.

Prestatie-indicatoren

Om de vraagvoorspelling te toetsen zijn allereerst prestatie-indicatoren nodig. Er zijn verschillende manieren om de nauwkeurigheid van een voorspelling te meten. In de literatuur is veel geschreven en gediscussieerd over uiteenlopende indicatoren. In bijlage 1 zijn de in de literatuur voorkomende indicatoren beschreven. Hier zijn ook de voor- en nadelen van de verschillende indicatoren verwoord. Ondanks dat er vele indicatoren beschikbaar zijn, is er geen absoluut antwoord over welke 'de beste' is. Het is daarom lastig om met één enkele maat een voorspelling te beoordelen. Daarom is er in goed overleg gekozen om de volgende prestatie-indicatoren te gebruiken: MSE, A-MAPE, CFE, NOSp en de PIS. Onderstaande tabel geeft een overzicht wat deze indicatoren precies betekenen. Voor definities en andere opmerkingen verwijs ik naar bijlage 1.

Tabel 1 – Gekozen prestatie-indicatoren

Prestatie-indicator	Afkorting	Beschrijving
Mean Squared Error	MSE	Geeft de gemiddelde gekwadrateerde afwijking van de voorspelling
Adjusted Mean Absolute Percentage Error	A-MAPE	Geeft de gemiddelde procentuele afwijking van de voorspelling t.o.v. de gemiddelde daadwerkelijke vraag
Cumulated Forecast Error	CFE	Geeft de cumulatieve afwijking van de voorspelling
Number of Shortages (percentage)	NOSp	Geeft aan hoeveel procent van de voorspelde maanden er een tekort was (voorraadbeheer buiten beschouwing gelaten)
Periods in Stock	PIS	Geeft aan hoeveel en hoelang artikelen in voorraad liggen/tekort komen aan het eind van de horizon (voorraadbeheer buiten beschouwing gelaten)

De eerste twee maten van de tabel geven harde cijfers over hoe de voorspelling afwijkt van de werkelijkheid. De MSE is een veelgebruikte prestatie-indicator en is als volgt gedefinieerd:

$$MSE = \frac{1}{n} \sum_{t=1}^N (F_t - X_t)^2$$

Waar F_t de voorspelling van de vraag in maand t is en X_t de daadwerkelijke vraag in die maand. Het verschil hier tussen wordt gekwadrateerd en vervolgens wordt het gemiddelde genomen door te delen door alle beschouwde periodes (in dit geval 12 maandvoorspellingen per jaar).

De A-MAPE is een aangepaste variant van het veelgebruikte MAPE. De nieuwe variant wordt gebruikt omdat in de definitie van MAPE het verschil tussen de voorspelling en de daadwerkelijke vraag gedeeld wordt door de daadwerkelijke vraag (zie bijlage 1). Het is hier echter veelvoorkomend dat de daadwerkelijke vraag in een maand 0 is en daarom is de MAPE niet gedefinieerd in die gevallen. De aangepaste variant deelt het gemiddelde verschil tussen de voorspelling en de vraag door de gemiddelde vraag. Hierdoor is de kans op delen door 0 niet meer aanwezig.

$$A - MAPE = \frac{\frac{\sum_{t=1}^n |F_t - X_t|}{N}}{\frac{\sum_{t=1}^n X_t}{N}}$$

De implicaties van de bovenstaande definitie is dat de A-MAPE minder vanzelfsprekend te interpreteren is dan de MAPE. Een waarde van 100% wil zeggen dat de gemiddelde afwijking precies eenmaal de gemiddelde vraag is. Aangezien het voor een voorspelling bijna onmogelijk is om de echte vraag precies te beschrijven, zijn ook hoge waarden van de A-MAPE te verwachten. Het is daarom niet te verwachten deze prestatie-indicator voorspelfouten zal geven van enkele procenten, maar eerder enkele tientallen procenten. De norm die hier gesteld wordt bedraagt 100%, dezelfde norm als in het werk van (Callegaro, 2010). Een zo laag mogelijke A-MAPE (30 á 40% is realistisch) is natuurlijk wenselijker.

De bovenstaande maten geven echter niet aan hoe groot of klein de gevolgen zijn van een verkeerde voorspelling. Daar zijn de laatste drie maten voor. De eerste van deze drie geeft de cumulatieve afwijking aan van de voorspelling aan het eind van horizon T (in dit geval één jaar). Een waarde dicht bij 0 moet aangeven dat de voorspelling juist is geweest. Wat 'dicht bij 0' definieert hangt af van de orde van grootte van de gegevens. Een waarde van 0 wil niet direct zeggen dat de voorspelling helemaal perfect is geweest aangezien eerdere onder-voorspelling gecompenseerd kan worden door latere over-voorspelling of vica versa. Om meer inzicht te geven wordt ook het maximum en het minimum van de cumulatieve afwijking bijgehouden. Een negatieve CFE duidt op een over-voorspelling en een positieve CFE duidt op een onder-voorspelling. Dit komt omdat in de definitie de voorspelde waarde van de echte waarde afgetrokken wordt. Zie onderstaande vergelijking.

$$CFE_T = \sum_{t=1}^T (X_t - F_t) = X_T - F_T + CFE_{t-1}$$

Om een beeld te geven hoe goed of slecht de voorspelling was in termen van het aantal periodes teveel/te weinig voorspeld is gekozen om de cumulatieve fout aan het einde van de horizon (CFE_T) te delen door de gemiddelde maandelijkse vraag ($\hat{X}_T$). Dit quotiënt geeft aan hoeveel maanden (gemiddelde) vraag er teveel of te weinig is voorspeld. Daarnaast is door deze deling de prestatie-indicator schaalonafhankelijk gemaakt en kan zo dus vergeleken worden tussen verschillende artikelen. Het quotiënt tussen CFE_T en $\hat{X}_T$ wordt hierna aangeduid met Cumulated Forecast Error in Periods (CFEp). Het teken wordt bij het quotiënt omgedraaid zodat een positieve waarde duidt op periodes teveel voorspeld en een negatieve waarde op periodes te weinig voorspeld.

Hoe erg een cumulatieve afwijking is gedurende de horizon, wordt weergegeven aan de hand van de Number of Shortages. Een shortage ontstaat als de cumulatieve afwijking in een bepaalde maand

groter is dan 0 (onder-voorspelling) en de vraag in die maand ook groter is dan 0. Dit is het best voor de geest te halen als je je beseft dat wanneer de cumulatieve afwijking positief is dat er de afgelopen periodes steeds te weinig is voorspeld. Als er steeds te weinig is voorspeld creëer je een fictief tekort. Als er dan vraag naar dat artikel komt, kun je niet leveren. Wanneer de cumulatieve afwijking positief is en er is geen vraag, is er dus ook geen sprake van een shortage. Dit is in de onderstaande formule weergegeven.

$$\text{Als } X_t \neq 0 \text{ \& } CFE_t > 0 \text{ dan } NOS \leftarrow NOS + 1$$

Er wordt hier gesproken over shortage en niet over stock-out, omdat voorraadbeheer buiten beschouwen gelaten wordt. Het kan best zijn dat met een enorme voorraad uit voorgaande jaren er nog steeds geleverd kan worden (ook al is je voorspelling systematisch te laag), maar dit is geen indicator voor hoe goed de voorspelling is. In dit onderzoek wordt de Number of Shortages uitgedrukt als een percentage over de totale horizon (NOSp genoemd).

Het kan best voorkomen dat over de complete horizon de NOSp gelijk is aan 0%. Dit wil dus zeggen dat je of perfect hebt voorspeld of systematisch hebt over-voorspeld. Hoe erg een onder- of over-voorspelling is, wordt gemeten aan de hand van de Periods in Stock (PIS). De PIS zegt hoeveel stuks er in voorraad liggen/tekort komen en hoelang deze al in voorraad liggen/tekort komen.

$$PIS_T = - \sum_{t=1}^T CFE_t$$

De definitie van de PIS aan het einde van de horizon (PIS_T) is dan ook het omgekeerde van alle cumulatieve afwijkingen, omdat een over-voorspelling negatieve CFE-waardes geeft. Als bijvoorbeeld over een horizon van twee maanden, in maand 1 vijf teveel werd voorspeld en in maand 2 tien te veel. Dan geldt $CFE_1 = -5$, $CFE_2 = -15$ en $PIS_2 = 20$. Deze waarde 20 wil in dit geval dus zeggen er 5 stuks 2 maanden lang in (fictieve) voorraad liggen en 10 stuks 1 maand lang.

De combinatie van NOSp en PIS zegt dus iets over het aantal tekortkomingen en de amplitude van deze tekortkomingen. Het kan namelijk best zo zijn dat de NOSp bijna 100% is, omdat er bijvoorbeeld systematisch net iets te weinig wordt voorspeld elke maand. Maar omdat de mate van tekortkoming in dit geval laag is, zal de PIS dicht bij 0 moeten liggen. Wat 'dicht bij 0' definieert hangt af van de orde van grootte van de gegevens.

Daarnaast wordt de prestatie van de voorspelling ten opzichte van de daadwerkelijke vraag gemeten aan de hand van lineaire regressie. De ideale voorspelling zou namelijk precies gelijk moeten lopen met de daadwerkelijke vraag. Dan zouden de cumulatieve vraag en de cumulatieve voorspelling zich 1 op 1 met elkaar verhouden. Wanneer je op de horizontale-as de cumulatieve vraag zou plotten en op de verticale-as de cumulatieve voorspelling, zou de meeste ideale lijn dus gelijk zijn aan $y=x^1$. Een richtingscoëfficiënt hoger dan 1 duidt op over-voorspelling en een richtingscoëfficiënt onder de 1 op onder-voorspelling. Omdat het veel ruimte kost om dit voor elk artikel te laten zien, beperken we ons tot de richtingscoëfficiënt en het startgetal van de lijn tussen vraag en voorspelling. Deze lijn is de meest ideale lijn die de datapunten beschrijft met behulp van de kleinste-kwadraten methode. Hoe goed deze lijn door de data loopt, wordt weergegeven met de regressiecoëfficiënt.

¹ Dit is natuurlijk geen realistische situatie omdat fluctuaties in de vraag nooit meteen worden opgevangen, maar op z'n vroegst een periode later.

Resultaten huidige voorspelling

Tabel 2 geeft een korte samenvatting van de prestatie-indicatoren. Alleen de prestatie-indicatoren die niet-schaalafhankelijk zijn, zijn weergegeven.

Tabel 2 – Prestatie-indicatoren huidige voorspelling

N=62	Gemiddelde	Std. Afwijking	Minimum	Maximum
A-MAPE	96,64%	51,04%	20,56%	286,67%
CFEp	2,63	6,858	-6,6	26,40
NOSp	39,80%	33,77%	0,00%	100,00%
R	0,9564	0,04032	0,8037	0,9985
Richtingscoëfficiënt	1,2301	0,5286	0,3688	3,0614

De gemiddelde A-MAPE bedroeg 96,64% met een minimum van 20,56% en een flinke uitschieter naar 286,67%. Zoals gezegd is de norm op 100% gesteld en gemiddeld gezien valt de huidige voorspelling voor de 62 artikelen er nog net onder. In zo'n 5 van de 12 maanden was er een tekort, voorraadbeheer buiten beschouwing gelaten (NOSp = 39,80%). De richtingscoëfficiënt blijkt juist op een over-voorspelling te duiden (gemiddelde 1,23), maar in ongeveer de helft van de gevallen bleek het startgetal negatief te zijn. Vandaar dat er gemiddeld (vooral in de eerste helft van het jaar) nog steeds veel shortages zijn. Daarnaast duidt een positieve CFEp ook op over-voorspelling aangezien er gemiddeld aan het einde van het jaar 2,63 maanden teveel werd voorspeld. Het merendeel van de artikelen had een positieve CFEp.

Al met al geeft de huidige voorspelling wisselende resultaten. Gemiddeld voldoet de A-MAPE wel aan de norm, maar er zijn ook flinke uitschieters naar boven. Ook de richtingscoëfficiënt zit gemiddeld niet in de buurt van de 1 en de standaardafwijking is ook nog eens groot (0,5286). Ook de Periods in Stock blijkt in ongeveer de helft van de gevallen positief te zijn en in de helft van de gevallen negatief, wederom duidend op wisselende resultaten. Daarnaast duidt de positieve CFEp op een voorspelling die gemiddeld voor een overschot zorgt aan het einde van het jaar.

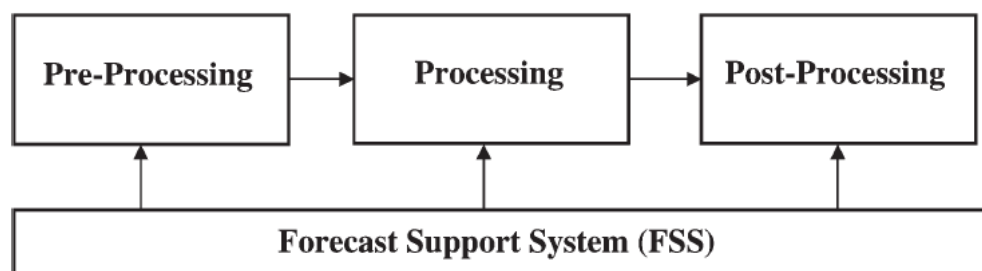
De huidige voorspelling houdt echter geen rekening met vraagkarakteristieken en gebruikt dezelfde methodiek (vierjaarlijks gemiddelde) voor alle soorten artikelen. In dit onderzoek wordt echter gekeken naar verschillende karakteristieken van vraagpatronen en daarbij horende classificaties en voorspelmethodieken. Aan de hand daarvan wordt deze data-set nog een keer geanalyseerd in termen van de prestatie-indicatoren. De uiteindelijk vergelijking tussen de huidige voorspelling en de nieuwe voorspelling wordt in hoofdstuk 7 gegeven.

Hoofdstuk 4 Spare parts classificatie

Dit hoofdstuk beschrijft eerst een model om vraagvoorspelling integraal te bekijken en vervolgens wordt een classificatie van spare parts voorgesteld.

Forecast Support System

Om het hele proces van vraagvoorspelling integraal te bekijken, maak ik gebruik van het Forecast Support System (FSS) opgesteld door (Boylan & Syntetos, 2010). Dit systeem bestaat uit drie stappen: pre-processing, processing & post-processing. Het eerste gedeelte bestaat uit het classificeren van spare parts ten behoeve van vraagvoorspelling. Het tweede gedeelte omvat het selecteren van de juiste voorspelmethode en het laatste gedeelte bevat eventuele wijzigingen om de voorspelling te verbeteren. Deze processen worden hieronder schematisch weergegeven.



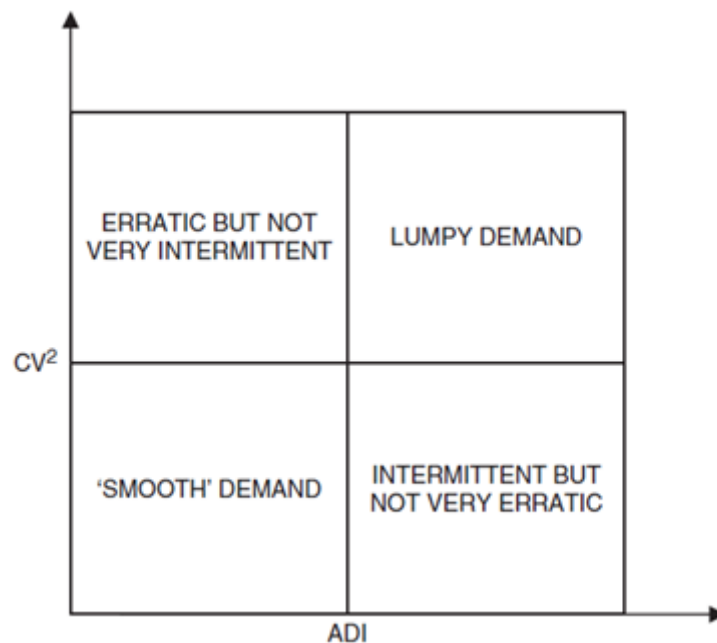
Figuur 2 – Fasen van voorspelling (uit Boylan & Syntetos, 2009)

Het pre-processing gedeelte wordt hieronder verder beschreven, het processing gedeelte wordt besproken in hoofdstuk 5, gevolgd door een data-analyse en empirische test in respectievelijk hoofdstuk 6 en 7. Het post-processing gedeelte komt in hoofdstuk 8 aan bod.

Theoretisch kader

Om de spare parts te classificeren op basis van vraagkarakteristieken is gezocht in de vakliteratuur naar een juiste benadering. In plaats van categorieën te bekijken die momenteel gangbaar zijn bij Stork (zoals productcategorie) worden deze allereerst genegeerd omdat geen enkele van deze categorieën impliciet relatie legt met vraagkarakteristieken. Omdat de volgende classificatie onafhankelijk gedaan is van eerdere inspanningen worden er geen aannames gemaakt of data bewust weg gefilterd. Daarnaast kan altijd later nog naar andere dwarsdoorsnedes gekeken worden die voor huidige classificatie-doeleinden geschikt zouden kunnen zijn. Op basis van een kort literatuuronderzoek is gekozen voor een tweedimensionale matrix, gedefinieerd door (Syntetos, Boylan, & Croston, 2005) en gebaseerd op (Williams, 1984). Dit theoretisch kader bleek veel naar gerefereerd te worden, door o.a (Boylan, Syntetos, & Karakostas, 2008), (Heinecke, Syntetos, & Wang, 2011), (Cavalieri, Garetti, Macchi, & Pinto, 2008) & (Wallström & Segerstedt, 2010). Ook een uitgebreid literatuuronderzoek naar verschillende SKU classificaties laat zien dat de classificatie ten behoeve van het selecteren van een geschikte voorspelmethode het beste gedaan kan worden op basis van vraagkarakteristieken door de tijd heen (van Kampen, Akkerman, & van Donk, 2012). Daarnaast is dit theoretisch kader ook uitvoerig getest in de praktijk. Door de twee dimensies is het model niet over ingewikkeld of lastig te interpreteren en daardoor leent dit model zich uitstekend om toegepast te worden in de praktijk. De uiteindelijke bedoeling van deze classificatie is om een geschikte voorspelmethode te kiezen met een superieure prestatie ten opzichte van andere methodieken. Welke methode superieur is per categorie wordt besproken in hoofdstuk 5.

In overleg met Stork is gekozen om te voorspellen op maandbasis. Dit vanwege de gemiddelde levertijden die het bedrijf kent alsmede de afzethoeveelheden per maand. De variabelen zijn zo gekozen dat ze rechtstreeks verband houden met de karakteristieken van de vraag. Deze twee variabelen zijn (1) de ADI: de gemiddelde tijd tussen vraagmomenten en (2) de CV^2 : de kwadratische variatiecoëfficiënt van de vraaghoeveelheid per maand. De eerste variabele geeft een maat voor hoe sporadisch vraag voorkomt en de tweede variabele geeft een maat voor de spreiding binnen vraagmomenten. Doordat apart wordt gekeken naar variatie en tussentijden kan een matrix worden opgesteld zoals te zien is in figuur 3.



Figuur 3 – Classificatie van vraagpatronen

In bovenstaand figuur staat ADI voor Average Demand Interval en CV^2 voor de kwadratische variatiecoëfficiënt van de vraaghoeveelheid per maand. De ADI is als volgt gedefinieerd:

$$ADI = \frac{1}{N-1} \sum_{n=2}^N (T_n - T_{n-1})$$

T_n = tijdstip van n^{de} maand met vraag,

Waar bij de ADI berekening de tijd tussen twee opeenvolgende maanden met vraag is gedeeld door het totaal aantal tussentijden ($N-1$). Als er bijvoorbeeld een artikel is met alleen vraag in maanden 1, 5 en 7. Dan is de ADI gelijk aan $\frac{1}{2} * ((5 - 1) + (7 - 5)) = 3$. Dit wil dus zeggen dat er gemiddeld drie maanden tussen vraagmomenten in zit.

De gekwadrateerde variatiecoëfficiënt is al volgt gedefinieerd:

$$CV^2 = \frac{\sigma^2}{\mu^2}$$

Met:

$$\sigma^2 = \frac{\sum_{i=1}^N (x_i - \bar{X})^2}{N}$$

En:

$$\mu = \frac{\sum_{i=1}^N x_i}{N}$$

x_i = vraaghoeveelheid in maand i ; $x_i > 0$

$\bar{X}$ = gemiddelde vraaghoeveelheid

N = totaal aantal maanden met vraag

Bij de variatiecoëfficiënt wordt de standaardafwijking gedeeld door het gemiddelde en vervolgens gekwadrateerd. Hierdoor krijg je een eenheidsloze maat die aangeeft hoe de standaardafwijking en het gemiddelde zich ten opzichte van elkaar verhouden. Deze berekening wordt uitgevoerd op basis van bestelhoeveelheidsgegevens wanneer er vraag is ($x_i > 0$). Perioden zonder vraag hebben dus geen invloed op het gemiddelde aangezien deze karakteristieken al worden meegenomen in de ADI. De ADI zegt dus iets over de spreiding tussen de vraagemomenten en de variatiecoëfficiënt zegt iets over de spreiding binnen vraagemomenten, beide geaggregeerd op maandbasis.

Wanneer een artikel een lage ADI heeft (dus een frequente vraag) en een lage variatiecoëfficiënt (weinig variatie ten opzichte van de gemiddelde vraaghoeveelheid) wordt deze geclassificeerd als smooth. Wanneer er producten zijn met een frequente vraag, maar wel erg variërend in hoeveelheid worden deze erratic genoemd. Wanneer een artikel een hoge gemiddelde tussentijd heeft, maar lage spreiding van de vraag dan wordt dit intermittent genoemd. Lumpy wordt benoemd als een artikel zowel een gemiddeld hoog interval heeft alsmede een hoge variatie binnen de vraag.

Scheidingswaarden

De matrix kent (tot nu toe) geen scheidingswaarden. De uiteindelijke bedoeling van de zojuist beschreven classificatie is om een superieure voorspelmethode te kiezen per categorie. Dit houdt dus in dat aan de hand van de ADI en CV^2 van een artikel een methodiek geselecteerd moet worden. De matrix dient gekwantificeerd te worden zodat een artikel met een bepaalde ADI en CV^2 ingedeeld kan worden in een kwadrant. Het volgende hoofdstuk beschrijft de methodieken die geschikt zijn per categorie en de bijbehorende scheidingswaarden.

Hoofdstuk 5 Voorspelmethodeken

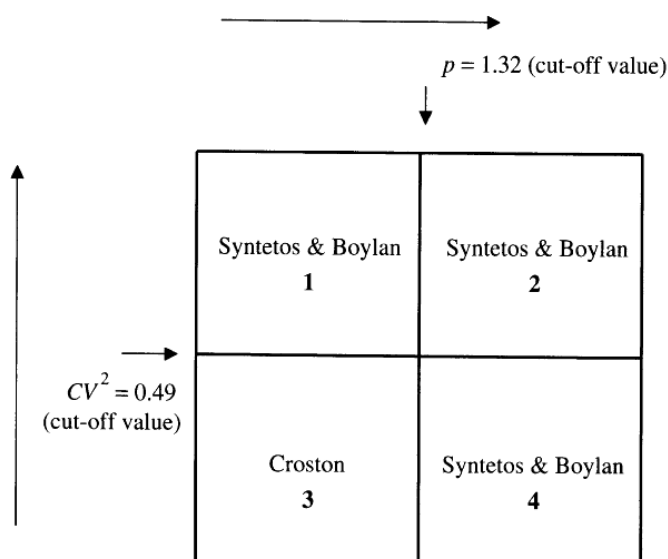
Dit hoofdstuk gaat in op de verschillende voorspelmethodeken die binnen de vakliteratuur worden behandeld. Dit hoofdstuk richt zich dus op de processing-fase van het FSS. Daarna wordt een selectie gemaakt welke methodeken geschikt zijn voor dit onderzoek.

Literatuuroverzicht

Om een beeld te krijgen welke methodeken er zijn en welke eventueel geschikt zouden zijn voor Stork, is in de literatuur gezocht. Deze zoektocht is alles behalve uitputtend en dient ook alleen om een algemeen beeld te krijgen van welke methodeken voorgeschreven worden en enkele karakteristieken te achterhalen per methodek. Nadat enkele methodes zijn geselecteerd kan dieper worden gekeken naar de materie in termen van exacte definities en implementatie. Bij deze zoektocht is dankbaar gebruikt gemaakt van (Callegaro, 2010), (Bacchetti & Saccani, 2012) en (Willemain, Smart, & Schwarz, 2004) die allemaal al verschillende methodeken in kaart hebben gebracht. Bijlage 2 geeft een overzicht van de in de literatuur gevonden methodeken met betrekking tot spare parts of intermittent/lumpy vraag.

Geschikte methodeken

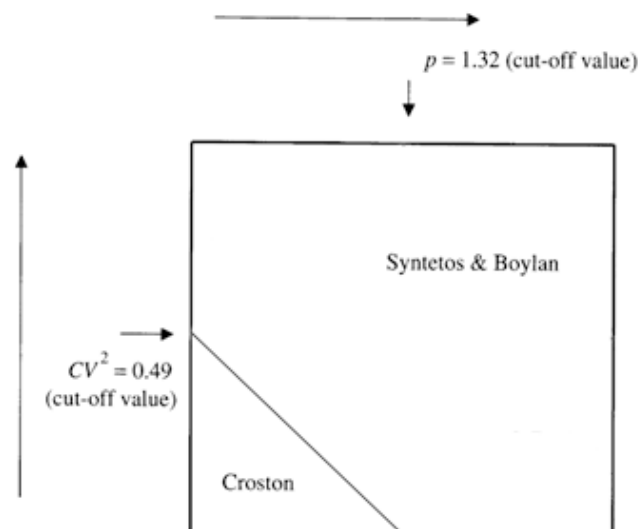
Het uitvoerig testen van de prestaties van alle gevonden methodeken zou simpelweg teveel tijd kosten en waarschijnlijk van weinig toegevoegde waarde zijn ten opzichte van een meer grove selectie. Daarnaast moet rekening gehouden worden met de praktischeerbaarheid van de methodeken. Het eerste interessante werk dat naar voren komt wordt beschreven in (Syntetos, Boylan, & Croston, 2005). Deze onderzoekers hebben scheidingswaarden gedefinieerd en op basis hiervan een indeling gemaakt van de best presterende voorspelmethodek. Zij hebben onderzocht en empirisch getest dat hun methodek (SB-Approximation) altijd beter werkt dan Croston (CR) mits de ADI (in het figuur p) of de CV^2 boven de scheidingswaarde is. Anders (in de categorie smooth dus) bleek CR beter te presteren dan SBA of de verschillen waren zo klein dat ze niet significant zijn. In figuur 4 staat de (theoretisch) best presterende voorspelmethodek per kwadrant van het in hoofdstuk 4 beschreven theoretisch kader.



Figuur 4 – Superieure voorspelmethodek per kwadrant volgens Syntetos & Boylan

De waarden gepresenteerd in figuur 4 zijn vaste scheidingswaarden voor het kiezen van een superieure methodiek (Syntetos, Boylan, & Croston, 2005). Deze waarden komen uit een grondige theoretische analyse. In deze analyse is gekeken naar de voorspelfout van de ene methodiek ten opzichte van de andere methodiek. Deze voorspelfout is in dat onderzoek gekoppeld aan de ADI en de CV^2 . Het heeft weinig toegevoegde waarde om uitgebreid in te gaan op waarom deze scheidingswaarden zijn zoals ze zijn, maar voor de geïnteresseerde lezer refereer ik naar (Syntetos, Boylan, & Croston, 2005).

Deze waarden zijn echter nog wel in twijfel getrokken door (Kostenko & Hyndman, 2006). Zij hebben aangetoond dat deze vaste scheidingswaarden slechts een benadering zijn. Zij suggeren om een lineair verband te gebruiken tussen de ADI en de CV^2 . Daardoor zou het schema uit figuur 4 worden opgedeeld in slechts twee vlakken in plaats van vier (zie figuur 5). (Heinecke, Syntetos, & Wang, 2011) hebben deze benadering empirisch getest en kwamen tot de conclusie dat de lineaire benadering inderdaad een betere voorspelaccuratie laat zien dan de vaste scheidingswaarden. Deze verbetering is consistent, maar procentueel erg weinig en ook nog eens afhankelijk van de te gebruiken dataset. Daarnaast zijn implicaties voor voorraadbeheer nog niet bewezen. In dit onderzoek worden de vaste cut-off values gebruikt omdat a) het kwadrant intuïtief is, omdat het impliciet relaties legt tussen de aard van de vraagkarakteristieken, b) het kwadrant makkelijk praktisch toepasbaar is en c) het lineaire verband nog niet uitvoerig getest is in meerdere studies.



Figuur 5 – Superieure voorspelmethode volgens Kostenko & Hyndman

(Syntetos, Boylan, & Croston, 2005) hebben echter lang niet alle voorspelmethodieken getoetst. Een ander onderzoek heeft 13 veelgebruikte voorspelmethodieken getoetst voor intermitterend demand en concludeerde dat exponential smoothing (SES) en Croston betere presteerden dan de rest (Ghobbar & Friend, 2003). Dit onderzoek heeft echter niet de SBA onderzocht. Beide onderzoeken gecombineerd zou je kunnen concluderen dat Croston beter presteert dan de overige en SBA weer beter dan Croston, maar dit is slechts een aanname. Daarnaast is ook de prestatie van ingewikkeldere methodieken zoals het Grey prediction model (GM) en verschillende Neural Networks (NNs) onderzocht. (Amin-Naseri & Tabar, 2008) laten bijvoorbeeld zien dat op basis van drie verschillende prestatie-indicatoren (A-MAPE, MASE & PB) een NN beter presteert dan

bijvoorbeeld SBA. Aan de andere kant is ook aangetoond dat een GM en een relatief eenvoudig NN weinig beter presteren (1 tot 2 %) dan een simpele Moving Average (MA) methodiek (Chen & Chen, 2009).

Ook (Callegaro, 2010) laat zien dat een NN geen goede voorspellingen geeft, mede omdat het echte voordeel van een NN pas behaald wordt bij een goed uitgebouwd en getraind model. Daarnaast wordt onderstreept dat methodieken gebaseerd zijn op trends niet goed presteren bij spare parts omdat er over het algemeen geen sprake is van seizoenspatronen of trends.

Op basis van voorgaande, praktiseerbaarheid en tijd is gekozen om grotendeels vast te houden aan het schema met de vaste scheidingswaarden. Vooral omdat deze keuze rechtstreeks verband heeft met de SKU classificatie die eerder is gemaakt. Om een tegenhanger te vinden voor bijvoorbeeld intermittent of lumpy kom je al snel terecht bij complexe methodieken zoals een Neural Network. Een simpel NN is eventueel nog te overwegen, maar zoals bovenstaand stuk heeft laten zien, is deze moeite niet gerechtvaardigd voor een (grote) toename in voorspelaccuratie. Toch zouden ingewikkeldere methodieken zoals bijvoorbeeld beschreven in (Chen, Chen, & Kuo, 2010) wél significante verbeteringen laten zien, maar vanwege ingewikkelde wiskundige berekeningen, praktiseerbaarheid en tijdsdruk is gekozen dat de extra moeite die dit met zich meebrengt niet opweegt tegen de baten.

Omdat in meerdere van de onderzochte studies de Moving Average (MA) vaak verassend uit de bus komt bij lumpy en intermittent, wordt deze methode (ook vanwege zijn eenvoud) naast SBA getoetst voor deze categorieën. In onderstaand schema worden de methoden die per categorie geanalyseerd gaan worden opgesomd.

Tabel 3 – Gekozen methodiek per categorie

Categorie	Methodiek 1	Methodiek 2
Smooth	Croston (CR) ²	-
Erratic	Syntetos-Boylan Approximation (SBA)	-
Lumpy	Syntetos-Boylan Approximation (SBA)	Moving Average (MA)
Intermittent	Syntetos-Boylan Approximation (SBA)	Moving Average (MA)

² De methode van Croston is hetzelfde als de Simple Exponential Smoothing methodiek voor ADI=1.

Gekozen methodieken

De gekozen methodieken uit de voorgaande sectie (CR, SBA & MA) worden hier per stuk toegelicht en gedefinieerd.

Croston's method (CR)

De methode van Croston werd voor het eerst beschreven in (Croston, 1972). Tientallen jaren later wordt deze methodiek nog steeds besproken in de literatuur en wordt inmiddels veel gebruikt in de praktijk en is geïmplementeerd in diverse softwarepakketen. Het is in feite een verbetering van Single Exponential Smoothing (SES) door ook rekening te houden met de tussentijd-intervallen als er geen vraag is. De methodiek is als volgt gedefinieerd:

$$Z_t = \alpha * X_{t-1} + (1 - \alpha) * Z_{t-1} \quad (5.1)^3$$

$$P_t = \alpha * G_{t-1} + (1 - \alpha) * P_{t-1} \quad (5.2)$$

De variabelen zijn als volgt gedefinieerd:

Z_t = de voorspelling van de van de vraag op tijdstip t

X_t = de werkelijke waarde van vraag op tijdstip t

P_t = de voorspelling van de tijd tussen vraagmomenten op tijdstip t

G_t = de werkelijke waarde van de tijd tussen vraagmomenten op tijdstip t

α = een smoothing constante ($0 \leq \alpha \leq 1$)

De voorspelling voor de vraag per periode op tijdstip t is dan gegeven door het quotiënt van vergelijkingen (5.1) en (5.2):

$$F_{t+1} = \frac{Z_t}{P_t} \quad (5.3)$$

De methodiek werkt als volgt: vanaf $t=1$ wordt gekeken naar X_t en G_t (hierbij wordt aangenomen dat Z_{t-1} en P_{t-1} precies de goede voorspellingen voor $t=1$ zijn, zodat $F_1 = X_1$). Wanneer er vraag is zodat $X_t > 0$ worden beide formules geüpdatet. Wanneer dit niet het geval is zullen Z_t en P_t dus gelijk zijn aan Z_{t-1} en P_{t-1} . De constante α zorgt voor de gradatie in hoeveel historische vraag meegenomen wordt ten opzichte van vraag die dichter bij het heden ligt. Deze constante is vrij te kiezen tussen 0 en 1, waar bij een waarde van 0 de methodiek zich uitsluitend op historische vraag baseert en bij $\alpha = 1$ zal de methodiek zich alleen richten het laatste vraagmoment.

³ Deze formule is feitelijk de methodiek van SES.

Syntetos & Boylan Approximation (SBA)

(Syntetos & Boylan, 2001) hebben laten zien dat de methode van Croston biased (systematisch fout) is. Daardoor is gemiddeld de voorspelling altijd te hoog. Zij hebben een correctieve factor ontwikkeld en is enkele jaren later uitvoerig getest (Syntetos & Boylan, 2005). De methodiek werkt exact hetzelfde als hierboven beschreven. De bias-gecorrigeerde formule is echter gedefinieerd als volgt:

$$F_{t+1} = \left(1 - \frac{\alpha}{2}\right) \frac{Z_t}{P_t} \quad (5.4)$$

Z_t en P_t worden bepaald aan de hand van vergelijkingen (5.1) en (5.2). Ondanks dat de SBA een verbetering is van de originele methode van Croston is in de vorige paragraaf toch besloten om Croston toe te passen bij de categorie smooth. Dit komt omdat de corrigerende factor afhangt van de grootte van de ADI. Hoe groter de ADI hoe beter SBA presteert ten opzichte van Croston. Meer informatie over waarom SBA beter is dan Croston is vinden in bijlage 3.

Moving Average (MA)

De MA is een relatief eenvoudige techniek die een gemiddelde berekent over de afgelopen n periodes. Hoe groter de waarde n hoe meer periodes de methode terugkijkt, maar is daardoor minder reactief op veranderingen. De methode is als volgt gedefinieerd:

$$MA(n) = (X_t + X_{t-1} + X_{t-2} + \dots + X_n)/n \quad (5.5)$$

De variabelen zijn als volgt gedefinieerd:

n = het aantal te beschouwen periodes

X_t = de werkelijke waarde van de vraag in periode t

De voorspelling voor de volgende periode wordt dan gegeven door:

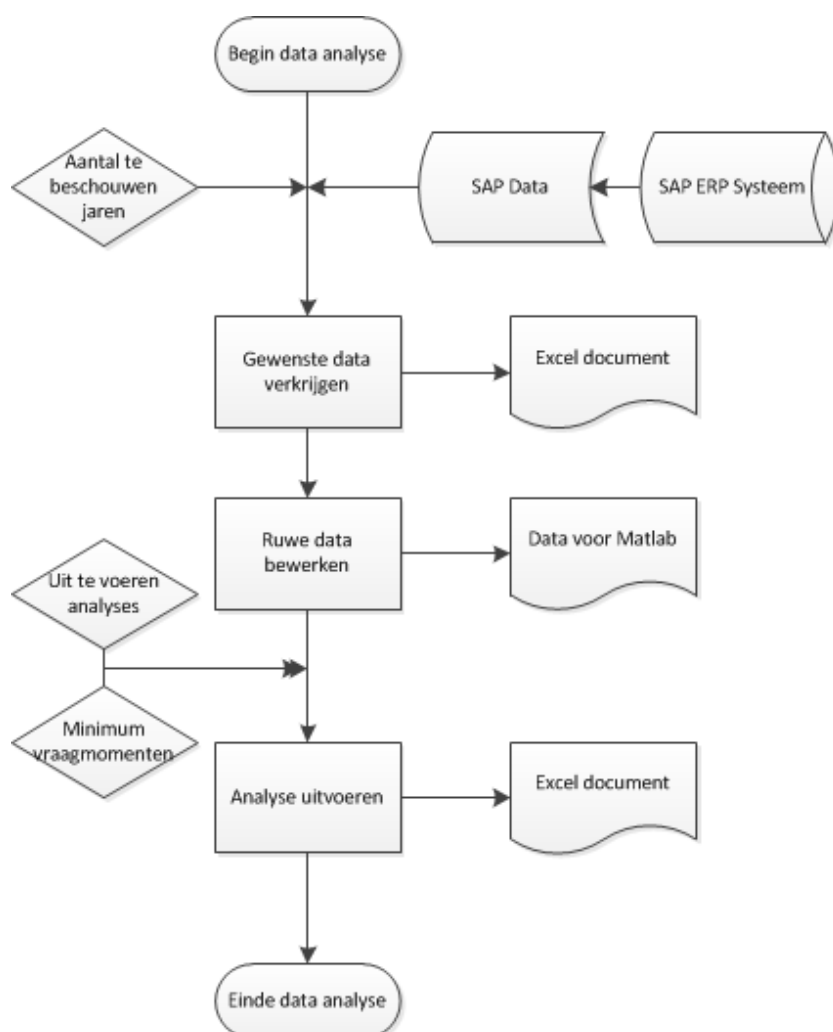
$$F_{t+1} = MA(n) \quad (5.6)$$

Hoofdstuk 6 Data analyse

Nu de pre-processing fase van het FSS gedefinieerd is in hoofdstuk 4, wordt in dit hoofdstuk deze fase gekwantificeerd aan de hand van een data-analyse.

Aanpak data analyse

Hiernaast is een schematische weergave gegeven voor de aanpak van de data analyse. Het SAP ERP systeem wordt gebruikt om de historische vraaggegevens op te vragen over een aantal jaar. In eerste instantie is gekozen voor een horizon van drie jaar. Daarna is de data bewerkt en vervolgens in Matlab geanalyseerd. De gegevens uit deze analyse werden vervolgens weer teruggekoppeld naar Excel. Dit proces is meerdere malen herhaald voor verschillende analyses. De meest gebruikte analyse is de invulling van het theoretisch kader door het berekenen van de ADI en de CV^2 per artikel. Meer gedetailleerde informatie over de data analyse is te vinden in bijlage 4.



Figuur 6 – Flowchart data-analyse

Resultaten data analyse

Tabel 4 – Aantal artikelen per categorie (data 2009-2011)

Minimum aantal vraagmomenten	Smooth	Erratic	Lumpy	Intermittent	Totaal
30	88	120	89	52	349
20	90	120	165	115	490
1	151	122	519	1956	2748

In tabel 4 is te zien hoeveel artikelen in elke categorie komen aan de hand van verschillende minimum aantal vraagmomenten. Er is goed te zien dat met het verlagen van het minimum er meer artikelen in de categorie lumpy en (vooral) intermittent komen. Als het minimum namelijk wordt verlaagd, neemt de kans op grotere tussentijden toe (en de ADI dus ook). Dit is enerzijds natuurlijk ook logisch, maar anderzijds moet je je afvragen of het zinvol om artikelen met zo weinig vraagmomenten te gaan voorspellen. Voor die artikelen is het logischer om bijvoorbeeld op basis van

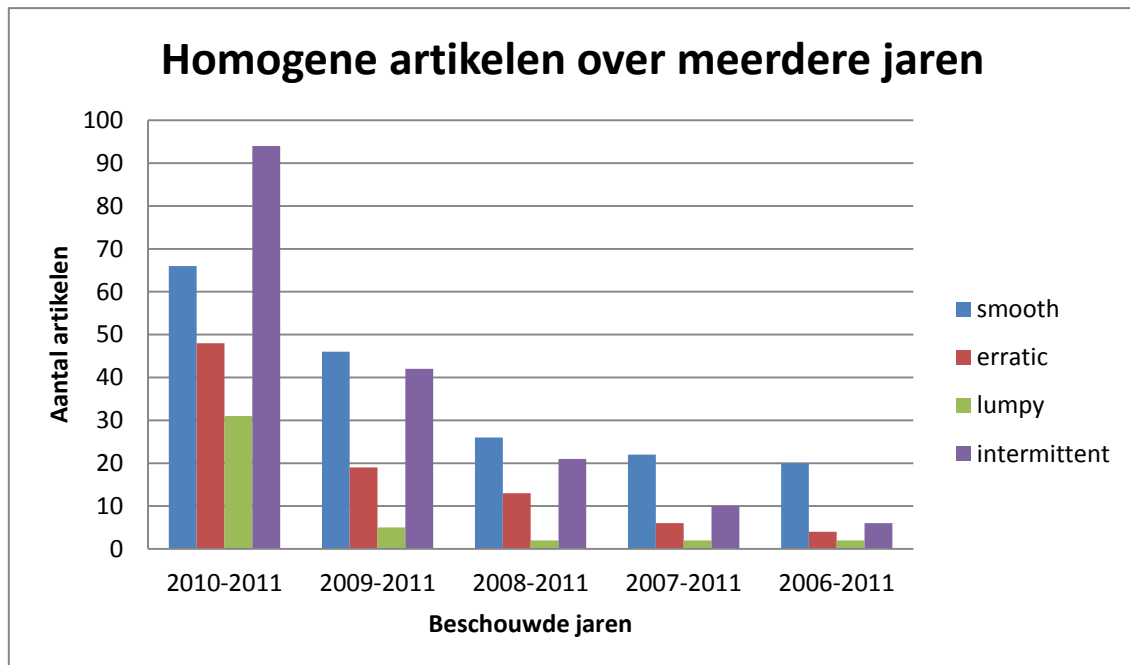
kritiekheid of levertijd enkele op voorraad te leggen. Daarnaast valt het aantal verkochte SKU's met meer dan 30 vraagmomenten erg tegen (349 ten opzichte van 4756). Het aantal artikelen dat geclassificeerd is onder de voorwaarde van minimaal één vraagmoment is 2748. Het verschil tussen dit aantal en de totaal verkochte hoeveelheid van 4756, zijn het aantal artikelen dat maar één keer verkocht is. Met één vraagmoment kan er namelijk geen ADI berekend worden. In bijlage 5 is het verband tussen de ADI & CV² en het minimum aantal vraagmomenten grafisch weergegeven.

Er zijn nog tal van andere analyses gedaan op de data. Zo is er gekeken of er coherente resultaten gevonden konden worden op een hogere aggregatieniveau zoals bijvoorbeeld productgroep, artikelgroep, of kritiekheid; deze bleven echter uit. Ook bleek dat het merendeel van de onderdelen door de jaren heen van categorie wisselde; dus het cumulatief bekijken van de data is geen representatieve weergave van de werkelijkheid. Toch wil Stork het liefst zo veel mogelijk onderdelen meenemen in de analyse zodat een groot aantal artikelen geclassificeerd en voorspeld kan worden. Dit is natuurlijk logisch, maar er moet ook goed gekeken worden naar een betrouwbare en goede weergave van de werkelijkheid. Daarom is in eerste instantie gekozen om over een bredere horizon te kijken (2006-2011) en hier de analyse over te doen. Daarnaast is een selectie gemaakt van artikelen die gedefinieerd zijn binnen een artikelgroep (zo'n 200 groepen zijn gedefinieerd); dit om alleen de artikelen mee te nemen die Stork wil zien in verdere analyse. Dit waren zo'n 3700 (verkochte) SKU's. Ook uit deze analyse bleek echter dat er weinig artikelen zijn met minimaal 20 of 30 vraagmomenten over zes jaar. Dit komt met name door grillige vraagpatronen, een wisselende productmix, een combinatie van nieuwe & oude machines (verschillende generaties) en andere veranderde aspecten die van invloed zijn op de vraagkarakteristieken zoals die gedefinieerd zijn.

Homogeniteitsanalyse

Om orde in de chaos te scheppen is in goed overleg besloten om te kijken naar onderdelen die niet van categorie wisselden over de laatste aantal jaren. Als een artikel hetzelfde geclassificeerd blijft noemen we dit een homogeen artikel. Uit deze analyse moet blijken bij welke artikelen de ADI en CV² niet veranderen. Daarbij opgemerkt dat het niet erg is als een onderdeel elk jaar erg gevarieerd is in hoeveelheden en elk jaar een lange gemiddelde tussenpoos kent (dus in lumpy valt), zolang de waarden van de twee variabelen maar dusdanig constant blijven door de jaren heen dat ze niet van categorie wisselen. De redenatie hierachter is dat deze artikelen sowieso aandacht verdienen voor voorspellingen en ook artikelen zijn die hoogstwaarschijnlijk komende jaren ook dezelfde karakteristieken kennen. Figuur 7 en tabel 5 illustreren het aantal artikelen dat door de jaren heen hetzelfde geclassificeerd bleven. De enige voorwaarde is een minimum aantal vraagmomenten van vijf per jaar. Met een hoger minimum bleken juist de lumpy en intermitterend artikelen weg te vallen, terwijl deze juist interessant zijn om te analyseren. Daarentegen zou een lager minimum, artikelen meenemen die zich niet lenen voor vraagvoorspelling. Met zo weinig vraagmomenten per jaar hangt de keuze om een artikel wel of niet op voorraad te leggen niet af van de vraag, maar van andere aspecten zoals prijs, levertijd en kritiekheid.

Het valt op dat over zes jaren (2006-2011) er slechts 32 artikelen zijn die constant blijven, d.w.z. en jaarlijks terugkomen en in dezelfde categorie. De analyse over de afgelopen twee jaren (2010-2011) laat een meer diverse invulling van de klassen zien: 239 artikelen. In deze twee jaren is het aantal SKU's dat meer dan 5 vraagmomenten kent en niet van klasse is veranderd zo'n 10% van het totaal aantal verkochte SKU's in 2011.



Figuur 7 – Aantal artikelen dat in de beschouwde jaren hetzelfde geclassificeerd bleven

Tabel 5 – Aantal artikelen dat in de beschouwde jaren hetzelfde geclassificeerd bleven

Categorie	2010-2011	2009-2011	2008-2011	2007-2011	2006-2011
Smooth	66	46	26	22	20
Erratic	48	19	13	6	4
Lumpy	31	5	2	2	2
Intermittent	94	42	21	10	6
Totaal	239	112	62	40	32

Zeer veel artikelen (~90%) zijn niet homogeen. Het komt vaak voor dat artikelen bijvoorbeeld eerst geclassificeerd waren als smooth, het jaar er na als erratic en vervolgens weer als smooth. Dit neigt naar de noodzaak voor een veel meer dynamischere blik op het gebied van vraagvoorspelling. Jaarlijks (zo niet maandelijks) kent elk artikel zijn eigen patroon.

In goed overleg is besloten om de artikelen die homogeen zijn in de jaren 2010-2011 te gebruiken voor verdere analyse. Hieronder staat een samenvatting van de gegevens over deze dataset. De gegevens zijn bepaald met behulp van IBM SPSS 20. Op pagina 21 zijn enkele voorbeelden te zien van het verschil in vraagpatroon tussen de verschillende classificaties.

Tabel 6 – Dataset voor verdere analyse (N=239)

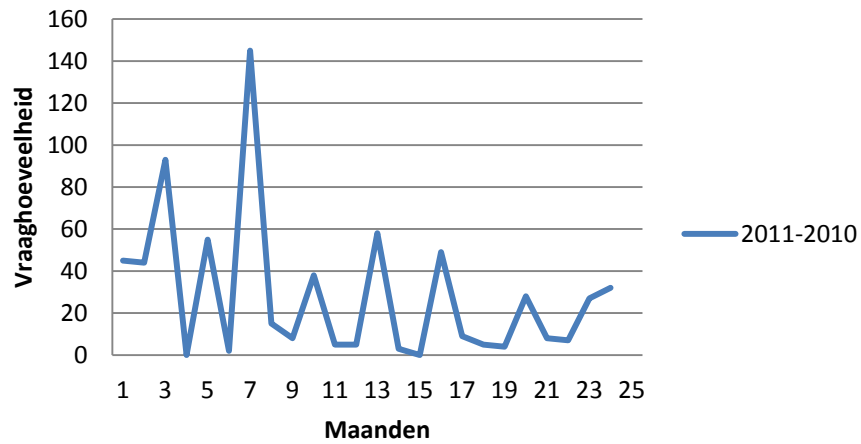
Classificatie	N	Percentage
Smooth	66	27.6%
Erratic	48	20.1%
Lumpy	31	13.0%
Intermittent	94	39.3%
Totaal	239	100%

Aanpak vervolgvraag

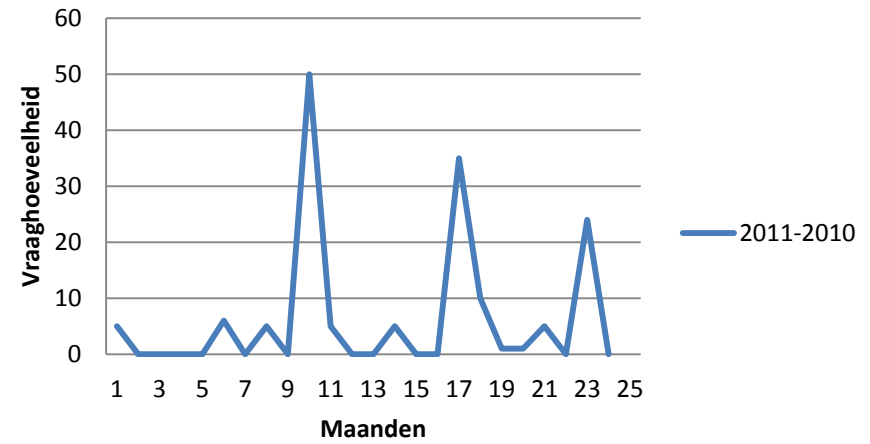
Ondanks het geringe aantal artikelen is de analyse op basis van homogeniteit wel geschikt om op een (statische) manier te onderzoeken welke voorspelmethodeken geschikt zijn en kunnen deze worden toegepast en getest op een echte (homogene) dataset. Deze zijn namelijk geschikt om een representatieve weergave van de werkelijkheid (en toekomst) te geven. Daarnaast bleken deze 239 stuks een hoge bijdrage te leveren aan de jaarlijkse omzet. Dankzij de homogeniteit zal de toetsing van een voorspelmethodek een betrouwbaar beeld geven en tevens is het inzicht in verschillende en geschikte voorspelmethodeken essentieel voor een eventueel latere (softwarematige) implementatie.

Het volgende hoofdstuk zal de reeds gekozen voorspelmethodeken per categorie gaan toetsen.

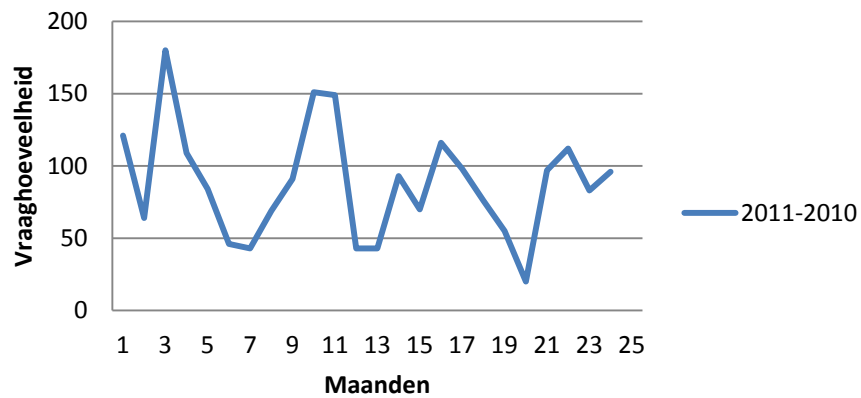
Erratic



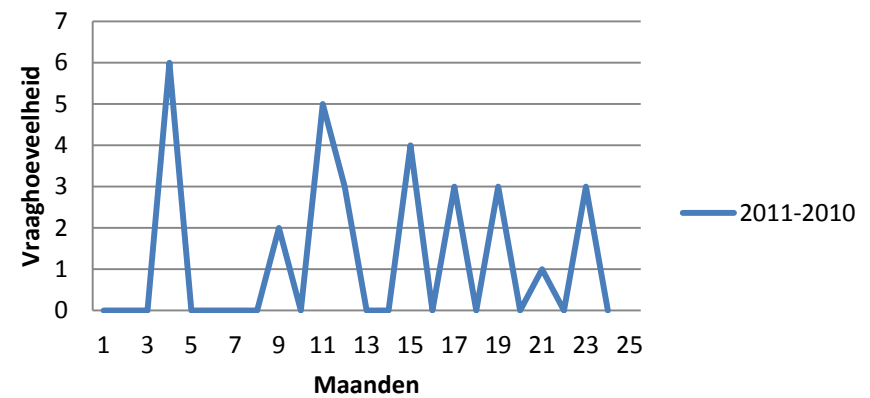
Lumpy



Smooth



Intermittent



Hoofdstuk 7 Empirische test

In dit hoofdstuk worden de gekozen methodieken getest aan de hand van prestatie-indicatoren.

Prestatie-indicatoren

De prestatie-indicatoren die gebruikt gaan worden zijn dezelfde zoals beschreven in hoofdstuk 3. Deze zijn daar gedefinieerd en beschreven en behoeven dus verder geen uitleg. De gekozen prestatie-indicatoren waren als volgt: MSE, A-MAPE, CFE, NOSP en de PIS. Daarnaast wordt aan de hand van lineaire regressie gekeken hoe goed de voorspelling zich verhoudt tot de daadwerkelijke vraag.

Test-set

Uit de data-set die aan het eind van hoofdstuk 6 is gedefinieerd, is een aantal artikelen gekozen om de voorspelmethodeken op toe te passen. Het analyseren van alle artikelen (239 stuks) uit die data-set zou simpelweg teveel tijd kosten en van weinig toegevoegde waarde zijn om te kijken of de voorspelmethodeken goed presteren. De gekozen artikelen zijn in overleg met de Product Manager gekozen op basis van twee criteria: 1) ze moeten uit de homogene data-set komen, en 2) de artikelen dragen substantieel bij aan de omzet. Een vuistregel die gehanteerd wordt binnen Stork is dat 20% van de spare parts voor zo'n 80% van de totale omzet aan spare parts zorgt. Daarnaast moest er een goede spreiding zijn over de categorieën.

Het aantal artikelen dat per categorie getoetst gaat worden, is weergegeven in onderstaande tabel.

Tabel 7 – Aantal te onderzoeken artikelen per categorie

Categorie	Aantal gekozen artikelen
Smooth	15
Erratic	17
Lumpy	15
Intermittent	15

Aanpak

De voorspelmethodeken worden toegepast in twee fases. Allereerst wordt er met behulp van de vraaggegevens over 2010 de parameters van het model geschat. Voor de SBA-methode is dit uitgebreid uitgewerkt in bijlage 6. Wanneer het model precies gedefinieerd is, wordt met dat model de vraag per maand per artikel van 2011 voorspeld (er steeds vanuit gaande dat de komende maand nog onbekend is). Aan het eind van elke maand wordt de nu wél bekende vraag van de afgelopen maand ingeladen. Zo worden 12 voorspellingen gedaan (januari t/m december 2011) en wordt aan het eind van het jaar gekeken hoe de voorspelling heeft gepresteerd. Dit kan gemakkelijk getoetst worden aangezien de echte vraag per maand van 2011 reeds bekend is.

Resultaten

Voor elke categorie is hier één voorbeeld uitgewerkt aan de hand van de prestatie-indicatoren. Daarnaast is in bijlage 7 voor elk voorbeeld een maandelijkse en een cumulatieve weergave van de voorspelling over 2011 gegeven. Ook een lineaire regressie wordt geplot in die bijlage. Een totaal overzicht met alle gegevens is opgenomen in bijlage 8.

Categorie: Smooth, Voorspelmethode: Croston

De eerste categorie is smooth en bevat artikelen met een lage CV^2 en een lage ADI. De voorspelmethode die hier het best bij past (zie hoofdstuk 5) is de methode van Croston. Onderstaande tabel sommeert de resultaten voor een enkel artikel.

Tabel 8 – Voorbeeld smooth: prestatie-indicatoren

MSE	A-MAPE	CFE	CFE max	CFE min	CFEp	NOSp	PIS	Lineaire regressie	
								Richtingscoëfficiënt	Startgetal
761	74%	-6	54	-11	0,2	58%	-124	1,136	-35

De A-MAPE duidt op een gemiddelde afwijking van 74% van de gemiddelde werkelijke vraag. In hoofdstuk 3 hebben we de norm al op 100% gezet, dit artikel valt daar dus redelijk onder. Daarnaast bleek de gemiddelde A-MAPE voor alle smooth-artikelen 55,93% te bedragen, een ruime verbetering van de norm dus. Aan het einde van het jaar was de cumulatieve fout nog -6 (dus 6 teveel voorspeld). Dit betekent een kleine over-voorspelling van ongeveer 0,2 maanden (CFEp). Daarbij wel opgemerkt dat gedurende het jaar het grootste tekort 54 stuks was (CFEmax) en het grootste cumulatieve overschot 11 (CFEmin). Dit duidt dus op hogere onder-voorspelling dan over-voorspelling gedurende het jaar. Dit is ook te zien aan het aantal tekorten. Er was een tekort in 7 van de 12 maanden (NOSp = 58%). De Periods in Stock bedroeg aan het eind van het jaar -124. Uit al deze gegevens kun je concluderen dat vooral in de eerste maanden van het jaar werd onder-voorspeld en dat gedurende het jaar dit iets werd goedge maakt door lichtjes over te voorspellen.

Ook uit de lineaire regressie blijkt dat er veelal wordt onder-voorspeld. Het richtingsgetal van dit artikel is weliswaar 1,136 (de cumulatieve voorspelling neemt meer toe dan de cumulatieve vraag), maar het startgetal bedraagt -35. Daardoor is de cumulatieve voorspelling ten opzichte van de vraag pas hoger na ongeveer de 10^{de} maand. De regressiecoëfficiënt van deze lijn door de data bleek ongeveer 0,984 te zijn. In bijlage 7 is een plot te vinden van de lineaire regressie.

Om uiteindelijk tot een zo hoog mogelijke beschikbaarheid te komen met een zo laag mogelijke voorraad zijn deze gegevens de sleutel tot succes. Als de richtingscoëfficiënt namelijk dicht bij de 1 is, volgt de voorspelling door het jaar heen aardig goed de vraag. Als het startgetal daarentegen negatief is, duidt dit op een onder-voorspelling. Een richtingscoëfficiënt dichtbij 1 is veel belangrijker dan een afwijkend startgetal (dit kan namelijk gemakkelijk gecorrigeert worden door (veiligheids)voorraad of een incidentele extra bestelling). De waarde 1 duidt namelijk op een consequente verhouding tussen voorspelling en vraag.

Tabel 9 geeft een samenvatting van alle resultaten (15 artikelen) in de categorie smooth. De tabel bevat alleen prestatie-indicatoren die niet schaalafhankelijk zijn, en dus vergeleken kunnen worden tussen verschillende artikelen.

Tabel 9 – Resultaten categorie smooth

Smooth (n=15)	Gemiddelde	Standaardafwijking	Minimum	Maximum
A-MAPE	55,93%	20,28%	17,75%	93,06%
CFEp	0,3	1,9482	-4,3	5,0
NOSp	57,77%	33,99%	8,33%	100,00%
R	0,98687	0,010190	0,967	0,999
Richtingscoëfficiënt	1,07913	0,200823	0,621	1,525

Uit deze tabel blijkt dat de methodiek van Croston toegepast in de categorie smooth voor aardig goede resultaten leidt. De A-MAPE ligt in alle 15 gevallen onder de 100%, met een gemiddelde van 55,93%. Toch blijkt uit de NOSp dat de methodiek wel systematisch onder-voorspelt (gemiddeld 58% shortages), maar dat deze onder-voorspelling wel consequent is (richtingscoëfficiënt ligt dichtbij de 1). Aan het einde van het jaar blijkt deze onder-voorspelling nog wel mee te vallen aangezien de CFEp gemiddeld slechts 0,3 maanden is. De voorspellingen zorgen aan het eind van het jaar dus zelfs nog lichtjes voor een overschot, maar is bijna gelijk aan 0. De voorspellingen zijn ook goed te analyseren aan de hand van deze richtingscoëfficiënt (en startgetal) aangezien de regressiecoëfficiënt (R) gemiddeld rond de 0,98 ligt.

Concluderend kun je zeggen dat de methode van Croston, ondanks dat hij iets onder-voorspelt, in z'n algemeenheid de werkelijke vraag aardig goed benaderd. Als een voorraadbeheerder de richtingscoëfficiënt en het startgetal van de voorspelling weet (en aangenomen dat de methode 2012 net zo goed voorspelt als 2011) kan hij zijn logistieke proces zo inrichten dat er de juiste hoeveelheid op het juiste tijdstip beschikbaar is.

Categorie: Erratic, voorspelmethode: SBA

In de categorie erratic (lage ADI, maar hoge CV²) is op basis van de literatuur gekozen voor de Syntetos & Boylan Approximation. Eén van de 17 geanalyseerde artikelen in de categorie erratic wordt hier toegelicht.

Tabel 10 – Voorbeeld erratic: prestatie-indicatoren

MSE	A-MAPE	CFE	CFE max	CFE min	CFEp	NOSp	PIS	Lineaire regressie	
								Richtingscoëfficiënt	Startgetal
102	79%	32	32	6	-3,3	92%	-169	0,8999	-8

Dit voorbeeld laat (net zoals de methode van Croston) een systematische onder-voorspelling zien. Dit komt met name omdat er nogal enkele uitschieters in de data zitten (een kenmerk van een erratic artikel vanwege de hoge variatie). Dit resulteert in een hogere A-MAPE, namelijk 79%. Maar dit blijft nog steeds onder de 100%. Ook de CFE, CFEmax en CFEmin duiden op een onder-voorspelling in elke maand (aan het eind van het jaar is er ongeveer 3 maanden vraag te weinig voorspeld). De NOSp en de PIS zijn daarom ook erg hoog. De gemiddelde maandelijkse vraag is zo'n 10 stuks en de PIS geeft dan ook een flinke negatieve waarde, dit duidt op serieuze onder-voorspelling in elke maand. Hoewel dit niet goed oogt, is er toch nog een positieve kant; de cumulatieve vraag lijkt een verschuiving te zijn ten opzichte van de daadwerkelijke vraag. Dit blijkt ook uit de regressieanalyse. In dit voorbeeld ligt de richtingscoëfficiënt op 0,8999 met een startgetal van -8. De onder-voorspelling houdt dus aan gedurende de hele horizon en verslechtert dus geleidelijk (richtingscoëfficiënt < 1). Dit is een goede weergave voor alle artikelen gezien de gemiddelde richtingscoëfficiënt 0,90 bedraagt en de startgetal zijn vaak negatief (hoewel relatief dicht bij de 0).

Tabel 11 – Resultaten categorie erratic

Erratic (n=17)	Gemiddelde	Standaardafwijking	Minimum	Maximum
A-MAPE	71,93%	9,52%	55,72%	95,20%
CFEp	-2,0	1,9932	-4,9	2,2
NOSp	78,43%	20,85%	33,33%	100,00%
R	0,97771	0,008372	0,964	0,988
Richtingscoëfficiënt	0,90300	0,20991	0,578	1,262

Bovenstaande tabel geeft de samenvatting van alle 17 onderzochte artikelen in de categorie erratic. In vergelijking met de categorie smooth hebben deze artikelen gemiddeld een hogere A-MAPE (71,9% tegen 55,93%), dit komt met name door de hogere variatie in de vraaghoeveelheden. Desondanks blijven alle artikelen nog netjes onder de 100%. Daarnaast is ook hier weer sprake van onder-voorspelling (alle artikelen hebben minimaal 4 van de 12 maand een tekort, NOSp-min = 33,33%). De CFEmax duidt op een onder-voorspelling van gemiddeld twee maanden vraag aan het eind van het jaar. Toch blijkt de richtingscoëfficiënt nog redelijk in de buurt te liggen bij de 1 (met waarden onder de 1 vaker voorkomend dan boven de 1). Ook alle startgetallen bleken negatief te zijn (twee uitzonderingen, beide rond de 0). Hieruit kun je dus concluderen dat de Syntetos & Boylan methode alle artikelen cumulatief gezien onder-voorspeld en dat deze onder-voorspelling steeds iets erger wordt door het jaar heen. Mochten deze resultaten representatief zijn voor alle artikelen die geclassificeerd kunnen worden als erratic is het belangrijk om voor een voorraadbeheerder goed

te kijken naar minimum voorraadhoogtes en minimum bestelhoeveelheden om toch nog aan de vraag te kunnen voldoen.

Categorie: Intermittent, voorspelmethode: SBA en MA

Artikelen in de categorie intermittent hebben een lage variatie, maar een hoge tussentijdinterval (CV^2 laag, ADI hoog). Hierdoor komt nul-vraag dus relatief vaak voor. De voorspelmethode SBA houdt hier rekening mee, maar de Moving Average (MA) niet. Het is dus allereerst noodzaak om een vergelijking te maken tussen de twee methodieken om uitsluitsel te geven welke methodek superieur is. Onderstaande tabel geeft het resultaat van de Wilcoxon Signed Rank Test om te onderzoeken of de prestatie-indicatoren significant van elkaar afwijken tussen de twee methodieken.

Tabel 12 – Resultaten Wilcoxon Signed Rank test voor SBA en MA in de categorie intermittent

Nulhypothese	Kans (p)	Beslissing	Significantieniveau (α)
De mediaan van het verschil tussen de MSE van SBA en de MSE van MA is gelijk aan 0.	0,917	Behoud de nulhypothese	0,05
De mediaan van het verschil tussen de A-MAPE van SBA en de A-MAPE van MA is gelijk aan 0.	0,328	Behoud de nulhypothese	0,05
De mediaan van het verschil tussen de PIS van SBA en de PIS van MA is gelijk aan 0.	0,311	Behoud de nulhypothese	0,05

Allereerst is gekeken of er significant verschil is tussen de MSE van beide methodieken, dit bleek absoluut niet waar te zijn ($p=0,917$, $\alpha=0,05$). Daarnaast zijn de verschillen tussen de A-MAPE en PIS van beide methodieken getest, maar ook deze gaven geen significant verschil aan ($p = 0,328$ en $p = 0,311$ respectievelijk, bij een gelijkblijvend significantieniveau). Omdat er geen significant verschil te vinden is tussen de twee methodieken⁴, wordt vastgehouden aan de oorspronkelijke aanbeveling uit de literatuur, namelijk het gebruiken van SBA in de categorie intermittent. Onderstaande tabel geeft weer een voorbeeld uit deze categorie.

Tabel 13 – Voorbeeld intermittent: prestatie-indicatoren

MSE	A-MAPE	CFE	CFE max	CFE min	CFEp	NOSp	PIS	Lineaire regressie	
								Richtingscoëfficiënt	Startgetal
11	113%	-9	9	-9	3,5	50%	-11	1,195	-5

Net zoals in de categorie smooth en erratic blijkt in dit voorbeeld sprake te zijn van onder-voorspelling. Dit is echter minder erg in de categorie intermittent door de vele momenten met nul-vraag. Omdat er geregeld nul-vraag is (5 van de 12 maanden in dit voorbeeld) zijn er dus ook minder shortages (gemiddeld is de NOSp zo'n 35%). Ook omdat er zo vaak nul-vraag is (en de methodek voorspelt gewoon positieve waarden) is de A-MAPE van 113% aan de hoge kant (ook het gemiddelde ligt boven de 100%). In dit voorbeeld wordt aan het eind van het jaar de schade van onder-

⁴ De hypothese is getoetst over een relatief laag aantal artikelen (15 stuks). Wellicht dat een grotere test-set andere resultaten geeft, maar voorlopig houden we vast aan de aanbeveling uit de literatuur.

voorspelling ingehaald door iets te veel te voorspellen (richtingscoëfficiënt = 1.195, startgetal = -5). Ook te zien aan de positieve CFEP aan het einde van het jaar. Dit blijkt ook de tendens voor de overige artikelen te zijn; de richtingscoëfficiënten zijn over het algemeen groter dan 1, maar de startgetallen zijn (licht) negatief. Daarnaast blijkt de data iets meer verspreid te liggen dan in de categorieën smooth en erratic (gemiddeld lagere R-waardes). Tabel 14 geeft een samenvatting van de 15 geanalyseerde artikelen in de categorie intermittent.

Tabel 14 – Resultaten categorie intermittent

Intermittent (n=15)	Gemiddelde	Standaardafwijking	Minimum	Maximum
A-MAPE	103,80%	22,76%	84,00%	153,85%
CFEP	1,8	4,7096	-3,3	12,0
NOSp	35,56%	26,63%	0,00%	75,00%
R	0,94680	0,025055	0,888	0,977
Richtingscoëfficiënt	1,18080	0,497199	0,504	2,357

De gemiddelde richtingscoëfficiënt en CFEP duidt op een over-voorspelling in deze categorie (ook in acht nemend de startgetallen die relatief dicht bij 0 liggen), de standaardafwijking bij beide indicatoren is echter meer dan een keer zo groot als bij de categorieën smooth en erratic. Het komt daarom voor dat sommige artikelen worden over-voorspeld terwijl een ander weer wordt onder-voorspeld. Voor alle artikelen geldt dat de PIS relatief dicht bij 0 ligt, dus er is geen sprake van systematische over- of onder-voorspelling.

Categorie: Lumpy, voorspelmethode: SBA en MA

Ook de categorie lumpy (hoge CV², hoge ADI) wordt geanalyseerd aan de hand van de SBA en de MA. Ook hier is weer een Wilcoxon Signed Rank Test gedaan om te kijken of de ene methode beter presteert dan de andere. Onderstaande tabel sommeert de resultaten.

Tabel 15 – Resultaten Wilcoxon Signed Rank test voor SBA en MA in de categorie lumpy

Nulhypothese	Kans (p)	Beslissing	Significantieniveau (α)
De mediaan van het verschil tussen de MSE van SBA en de MSE van MA is gelijk aan 0.	0,023	Verwerp de nulhypothese	0,05
De mediaan van het verschil tussen de A-MAPE van SBA en de A-MAPE van MA is gelijk aan 0.	0,005	Verwerp de nulhypothese	0,05
De mediaan van het verschil tussen de PIS van SBA en de PIS van MA is gelijk aan 0.	0,156	Behoud de nulhypothese	0,05

Opvallend is dat in deze categorie er wel verschillen zijn tussen de methodieken. De MSE en de A-MAPE wijken significant af ($p = 0,023$ en $p = 0,005$ respectievelijk). De Syntetos & Boylan Approximation presteert in deze categorie dus beter omdat de gemiddelde MSE en A-MAPE significant lager zijn dan bij de Moving Average. Tussen de PIS van beide methodieken is geen significant verschil gevonden. Daarom wordt de SBA toegepast in de categorie lumpy, hieronder wordt daar een voorbeeld van uitgewerkt.

Tabel 16 – Voorbeeld lumpy: prestatie-indicatoren

MSE	A-MAPE	CFE	CFE max	CFE min	CFEp	NOSp	PIS	Lineaire regressie	
								Richtingscoëfficiënt	Startgetal
19	111%	6	7	-6	-2,1	25%	-1	1,058	-1

Dit is een typisch voorbeeld van een lumpy-artikel: de vraag per maand varieert flink en bijna de helft van de maanden heeft geen vraag. De hoogste A-MAPE van alle categorieën is hier dan ook te vinden (gemiddeld zo'n 130% en het minimum komt niet onder de 90%). Toch blijkt in tegenstelling tot de andere categorieën de lumpy artikelen iets over-voorspeld worden. De NOSp is dan ook laag (gemiddeld 21,11%, ook dankzij de vele nul-vraag) en de PIS van de artikelen zijn over het algemeen positief (10 van de 15 artikelen). De richtingscoëfficiënt ligt dan ook gemiddeld boven 1, maar dan moet je wel genoeg nemen met een lagere regressiecoëfficiënt (gemiddeld $R=0,89447$), dit komt omdat door de hoge variatie én de vele nul-vraag van deze artikelen, de voorspelling ten opzichte van de vraag ook varieert.

Ondanks de relatief hoge afwijking van de voorspelling lijken de consequenties nog mee te vallen (in termen van NOSp en PIS) aangezien er sprake is van lichte over-voorspelling en nul-vraag. Dit betekent wel deze voorspelling voorraad opdrijvend kan worden. De resultaten van alle lumpy-artikelen staan hieronder in tabel 17.

Tabel 17 – Resultaten categorie lumpy

Lumpy (n=15)	Gemiddelde	Standaardafwijking	Minimum	Maximum
A-MAPE	127,18%	27,53%	90,20%	178,57%
CFEp	1,2	3,7948	-7,1	8,9
NOSp	21,11%	16,02%	0,00%	58,33%
R	0,89447	0,049578	0,799	0,953
Richtingscoëfficiënt	1,02520	0,392234	0,357	1,845

Zoals verwacht is de categorie lumpy het lastigst te voorspellen. Door de vele nul-vraag en de hoge variatie wil een voorspelling al snel uit de pas lopen. In het gegeven voorbeeld is er een flinke uitschieter in de laatste maand waar niet op wordt geanticipeerd (voorspelling = 2 stuks, werkelijke vraag = 14 stuks). Als het een goedkoop artikel is waar toch al veel van op voorraad lag, wil dit geen probleem zijn. Voor duurdere of kritische artikelen is het echter wel noodzakelijk om de vraag goed te kunnen voorspellen. Vooral in deze categorie geeft de voorspelmethode minder goede resultaten en zal er misschien naar andere manieren moeten worden gezocht om de resultaten te verbeteren. In plaats van ingewikkeldere methodieken te gebruiken wordt in het volgende hoofdstuk gekeken of de voorspelling 'aangepast' kan worden aan de hand van andere informatie. Dit zal dus de laatste fase van het Forecasting Support System beslaan; namelijk post-processing.

Onderstaande tabel sommeert de resultaten van alle categorieën. Het is duidelijk te zien dat de A-MAPE toeneemt naarmate de categorie een hogere CV² en/of ADI krijgt. De Number of Shortages neemt in de categorieën met een hoge ADI juist af, omdat met vele nul-vraag er ook geen tekort kan zijn. In de categorie smooth en erratic blijkt vooral onder-voorspeld te worden (richtingscoëfficiënten onder de 1), terwijl de categorie intermittent wisselende resultaten geeft en de categorie lumpy veelal wordt over-voorspeld.

Tabel 18 – Gemiddelde waarden van de prestatie-indicatoren voor alle categorieën

Categorie	A-MAPE	NOSp	CFEp	R	Richtingscoëfficiënt
Smooth	55,93%	57,77%	0,3	0,98687	0,98687
Erratic	71,93%	78,43%	-2,0	0,97771	0,90300
Intermittent	103,80%	35,56%%	1,8	0,94680	1,18080
Lumpy	127,18%	21,11%	1,2	0,89447	1,02520

Vergelijking huidige voorspelling met voorspelmethodeken

In het begin van dit onderzoek is gekeken naar de prestaties van de huidige manier van vraagvoorspelling. Om een eerlijke vergelijking te maken zijn alle prestatie-indicatoren gelijk getrokken. Allereerst moet onderzocht worden of de voorspelmethodeken betere resultaten geven ten opzichte van de huidige manier van voorspellen. Hiervoor wordt gebruikt gemaakt van de Wilcoxon Signed Rank Test voor alle 62 onderzochte artikelen.

Tabel 19 – Resultaten Wilcoxon Signed Rank Test voor huidige voorspelling met nieuwe voorspelling

Nulhypothese	Kans (p)	Beslissing	Significantieniveau (α)
De mediaan van het verschil tussen de MSE van de huidige voorspelling en de nieuwe voorspelling is gelijk aan 0.	0,006	Verwerp de nulhypothese	0,05
De mediaan van het verschil tussen de A-MAPE van de huidige voorspelling en de nieuwe voorspelling is gelijk aan 0.	0,500	Behoud de nulhypothese	0,05
De mediaan van het verschil tussen de PIS van de huidige voorspelling en de nieuwe voorspelling is gelijk aan 0.	0,267	Behoud de nulhypothese	0,05

Allereerst is de Mean Squared Error (MSE) tussen beide voorspellingen onderzocht. Hieruit bleek dat juist de huidige voorspelling betere resultaten geeft ($p=0,006$, $\alpha = 0,05$). Dit komt met name omdat de onder-voorspelling van de methodeken iets groter is dan de over-voorspelling van de huidige manier. Omdat de MSE absolute waarden vergelijkt, komt hier de huidige manier van voorspellen (dus een vierjaarlijkse gemiddelde, uniform over het jaar verdeeld) beter uit de bus. Aan de andere kant is er geen significant verschil tussen de A-MAPE en de PIS van beide methodeken. De A-MAPE van de voorspelmethodeken doet het gemiddeld iets beter (90% tegen 96%), maar al met al wijkt dit niet significant af. Het absolute verschil tussen de beide Periods in Stocks (dus te weinig leveren weegt net zo zwaar als teveel op voorraad) blijkt ook niet significant af te wijken.

Dit lijkt misschien een domper, maar er moeten wel enkele kanttekeningen worden bijgeplaatst. Zo is een vierjaarlijkse gemiddelde niet erg reactief op veranderingen en is het ongeschikt voor nieuwe of uitfaserende producten. Daarnaast zijn de implicaties voor het voorraadbeheer niet bekend. Omdat de nieuwe voorspelmethodeken beter met de dynamieken van de vraag omgaan, is het aannemelijk dat de consequenties voor het voorraadbeheer beter zullen zijn. Omdat de voorspelmethodeken met grotere fluctuaties voorspellen dan een uniforme verdeling (en daardoor een verminderde voorspelaccuratie) kan het best betekenen dat door die fluctuaties de voorraadhoogtes en het aantal nee-verkopen juist beter of elkaar wordt afgestemd. Zo zal een uniforme benadering geen reactie hebben op een flinke piek aan het begin van het jaar, terwijl de voorspelmethodeken daar wel op kan reageren. Hierdoor kan het zo zijn dat een tekort in de eerste maanden eerder wordt weggewerkt met de voorspelmethodeken dan met de huidige manier van werken. In hoofdstuk 9 wordt hier een voorzichtige blik op geworden.

Om de voorspelling op artikelniveau slim aan te pakken dient gezocht te worden naar een dynamische softwarematige tool die de voorspelmethodeken nog beter weet toe te passen (door bijvoorbeeld uitschieters er uit te filteren, patronen & trends te herkennen, etc.) en dit combineert met andere inzichten en kennis die van toegevoegde waarde kunnen zijn (bijvoorbeeld het betrekken van niet-historische data). Dit wordt besproken in het volgende hoofdstuk.

Hoofdstuk 8 Het aanpassen van de voorspelling

Dit hoofdstuk richt zich op de post-processing fase van het FSS door te kijken of voorspellingen verbeterd kunnen worden aan de hand van andere informatie dan historische data. Allereerst wordt gekeken hoe dit zou kunnen en daarna wordt een test gedaan om te kijken of een voorspelling verbeterd kan worden.

Aanpassingen aan de voorspelling

De voorspelling per maand per artikel kan om verschillende redenen niet tot goede resultaten leiden. Als verondersteld wordt dat het voorraadbeheer optimaal wordt ingericht en er nog steeds hoge voorraden zijn of veel tekorten, dan kan dit komen door een niet-optimale voorspelling. De voorspelling is tenslotte alleen gebaseerd op historische gegevens en niet op veranderingen in de markt, wisselingen in productaanbod, nieuwe generaties machines en andere variabelen. Daarom kan het nuttig zijn om niet alleen naar de statistische analyse te kijken, maar ook naar andere informatie die relevant kan zijn voor de vraaghoeveelheden. Zo kan Stork bijvoorbeeld op afstand tellers uitlezen op de machines van de klant en zo kan inschatten wanneer er een onderhoudsbeurt gedaan moet worden. Geïntensifieerd klantcontact geeft daarom ook meer inzicht in de behoefte van de klant en dus wellicht ook in zijn of haar bestelmomenten. Het is niet lastig om te bedenken dat alle bestellingen die je van te voren weet (deterministisch) niet voorspeld hoeven te worden.

In de ideale wereld zou Stork alle bestellingen van alle klanten van te voren weten (door bijvoorbeeld klantcontact, tellers uitlezen of andere informatiebronnen) en is de voorspelling per maand gelijk aan de bestelhoeveelheden die Stork al weet. Als bijvoorbeeld aan de hand van deze informatie enkele uitschieters van te voren bekend waren, zouden deze uitschieters uit de 'normale' voorspelling gehaald kunnen worden. Daardoor wordt de vraag meer regelmatig en dus ook beter te voorspellen. Alle bestellingen die je van te voren al weet, hoeven dan niet meer voorspeld te worden. Zo krijg je een voorspelling die voor een deel gebaseerd is op alles wat je al weet en voor een deel op historische data. Daarnaast kan een voorspelling verbeterd worden aan de hand van de expertise van een manager die ervaring heeft op dit gebied en bijvoorbeeld vindt dat de voorspelling aan de lage of hoge kant is (Syntetos A. A., Nikolopoulos, Boylan, Fildes, & Goodwin, 2009). Ook zou je bij de maandelijkse update van de voorspelling kunnen kijken, in het geval van een uitschieter, of deze wel of niet meegenomen wordt bij het berekenen van de behoefte van de volgende maand. Zo zijn er tal van mogelijkheden om extra kennis in de voorspelling te brengen. Denk bijvoorbeeld aan uitfasering van producten, nieuwe producten of revisies aan bestaande producten. In dit onderzoek richten we ons op het intensifiëren van klantcontact om zo beter in te schatten wanneer een klant behoefte heeft aan spare parts.

Geïntensifieerd klantcontact

Momenteel is Stork nog niet zo ver om van al hun klanten te weten wanneer ze bestellingen gaan plaatsen. Bij enkele klanten is dit al wel het geval en Stork wil ook meer aandacht vestigen op klantcontact om zo niet alleen de klant tevreden te houden, maar ook om waardevolle informatie uit te wisselen. Relevante informatie voor dit onderzoek zou bijvoorbeeld zijn om te weten wanneer een klant wat gaat bestellen. Dit is voor het gros van de gevallen nu nog niet zo, maar alles wat je al weet hoef je in ieder geval niet te voorspellen. Dus een iets meer deterministische vraag zou de voorspelling misschien al enorm ten goede komen. We nemen de proef op de som door te veronderstellen dat Stork heeft geïnvesteerd in verbeterd klantcontact en dat daardoor bij sommige klanten goede afspraken zijn gemaakt over het (ver) van te voren aangegeven wat ze willen

bestellen. De volgende paragraaf kijkt met behulp van een kleine test of deze informatie van toevoegde waarde is voor de kwaliteit van de voorspelling.

Orders inleggen

Om te kijken hoe bekende vraag de voorspelling kan verbeteren, wordt verondersteld dat enkele orders van te voren bekend zijn. De aanpassing aan de voorspelling moet echter wel met wijsheid gedaan worden. Er moet allereerst gekeken worden of de ingegde orders per maand de totaal te verwachten orders evenaren of dat verwacht wordt dat er nog meer van een artikel besteld gaat worden. Als bijvoorbeeld het gros van de klanten hun bestelling van te voren heeft doorgegeven, kun je met hoge(re) zekerheid zeggen dat dit de lading dekt. Dan zou de totaalhoeveelheid van de ingegde orders gelijk zijn aan de voorspelling per maand. Als er aan de andere kant echter verwacht wordt dat er meer orders komen voor hetzelfde artikel dan bekend, is het slim om wat extra te voorspellen. Dit kan bijvoorbeeld door de voorspelling (of een deel daarvan) bij de ingegde orders op te tellen. Het is hier al duidelijk dat door de aanpassing, de voorspelling subjectiever wordt. Er moet onderstreept worden dat degene verantwoordelijk voor de vraagvoorspelling met expertise en ervaring moet beslissen om een voorspelling al dan niet aan te passen aan de hand van bekende orders (of andere inzichten). Omdat het nu niet onderzocht kan worden hoe de betreffende persoon zou gaan handelen, is gekozen om een simplistisch voorbeeld te nemen.

Als er in een maand een order is ingelegd, wordt deze order van de daadwerkelijke vraag afgetrokken (zodat de voorspelling voor de volgende periode alleen afhangt van de niet-bekende vraag).

Vervolgens komt de keuze om vast te houden aan de voorspelling, de ingegde orders als voorspelling te zien of een mix van de ingelegde orders en de voorspelling. In dit voorbeeld hebben we voor twee (eenvoudige) beslisstrategieën gekozen. Strategie 1 is als volgt: als de voorspelde waarde groter is dan de ingegde order, wordt de voorspelling als leidraad genomen. Als de ingegde orders groter zijn, wordt deze als voorspelling gekozen (er wordt dus verondersteld dat de volledige vraag die maand dan bekend is). Strategie 2 kiest ook de voorspelde waarde als deze groter is dan de orders, maar telt de voorspelling en de ingelegde orders bij elkaar op als de orders groter zijn dan de voorspelling (er wordt dus verondersteld dat de ingelegde orders geen totaal beeld geven van de vraag per maand). Het is hier niet belangrijk om een uitgebreide 'wat-als'-analyse te doen, maar wel om het belang van deterministische informatie in het voorspelproces te onderstrepen. Daarom hieronder een illustratief voorbeeld.

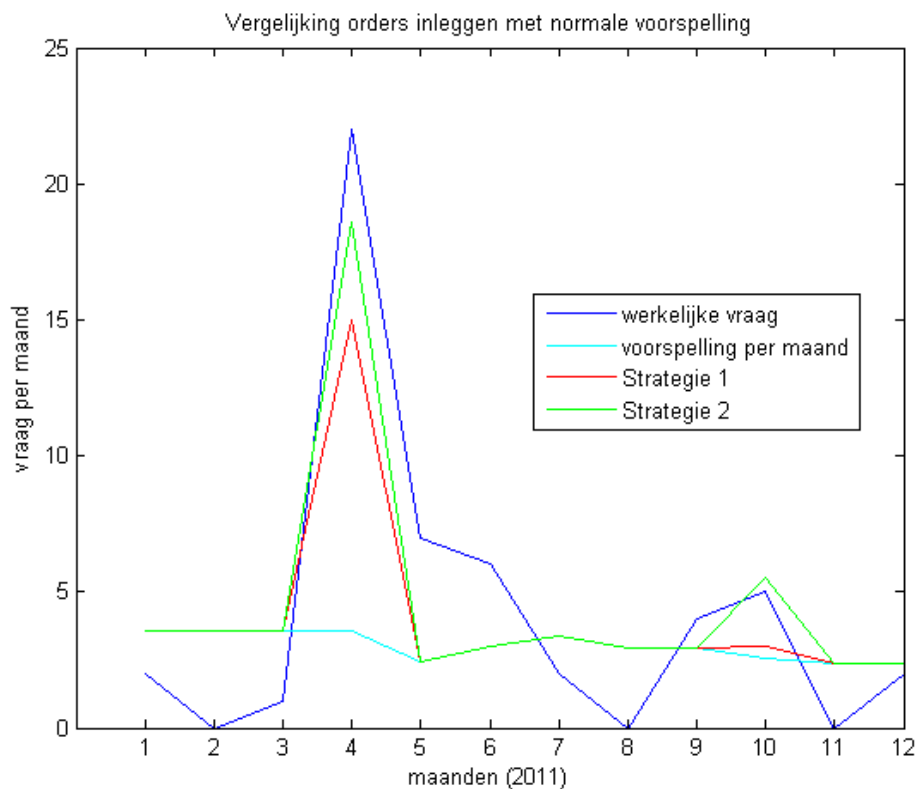
Voorbeeld orders inleggen

Onderstaande tabel geeft weer hoeveel vraag we van te voren hebben ingelegd bij een willekeurig gekozen artikel in de categorie lumpy. In vijf van de twaalf maanden is dus een deel van de vraag al van te voren bekend.

Tabel 20 – Bekende vraaghoeveelheid per maand

Maand	1	2	3	4	5	6	7	8	9	10	11	12
Vraag	2	-	1	15	-	2	-	-	-	4	-	-

Onderstaand figuur laat het verschil zien tussen de verschillende voorspellingen. De lichtblauwe lijn laat de voorspelling per maand zien (zonder dat er aanpassingen gedaan zijn), de rode lijn beschrijft strategie 1 en de groene lijn geeft strategie 2 weer. De donkerblauwe lijn is zoals altijd de werkelijke vraag. Er is te zien dat beide strategieën de voorspelling verbeteren. In slechts twee van de twaalf maanden (maanden 4 en 10) bleken de ingeleegde orders groter te zijn dan de voorspelling die al gedaan was. Er zijn dus eigenlijk in maar twee maanden echte aanpassingen gedaan op de voorspelling⁵.



Figuur 8 – Resultaten orders inleggen

Het inleggen van de orders heeft de prestatie van de voorspelling stukken verbeterd. Onderstaande tabel laat het verschil zien tussen de prestatie-indicatoren van de oorspronkelijke voorspelling (zonder ingeleegde orders) en tussen de beide strategieën.

Tabel 21 – Vergelijking tussen de normale voorspelling en de aangepaste voorspelling

	MSE	A-MAPE	CFE	CFEmax	CFEmin	NOSp	PIS
Oorspronkelijke voorspelling	38	90,20%	-4	11	-9	50%	-21
Strategie 1	10	66,67%	-2	4	-9	33,33%	13
Strategie 2	7	58,86%	-9	0	-9	0	58

Op een slimme manier de voorspelling aanpassen met behulp van extra informatie (zoals deterministische vraag) kan dus goede resultaten boeken. In de praktijk moet echter goed gekeken worden of de aanpassing gerechtvaardigd is zodat de voorspelling in ieder geval niet verslechtert.

⁵ De 'normale' voorspelling per maand is hier gebaseerd op de werkelijke vraag min de ingeleegde orders. De orders hebben dus indirect wel invloed op de voorspelling, ook al wordt de voorspelling niet aangepast.

Hoofdstuk 9 Consequenties voor voorraadbeheer

Dit hoofdstuk gaat kort in op de extensie van dit onderzoek, namelijk voorraadbeheer. Er wordt gekeken wat de interactie is van de voorspelmethodieken met het voorraadbeheer.

Aan het eind van hoofdstuk 7 hebben we geconcludeerd dat in termen van voorspelaccuratie (MSE, A-MAPE, PIS) de nieuwe voorspelmethodieken niet significant beter presteren dan de huidige voorspelmethode. We hebben daar ook beredeneerd dat de nieuwe voorspelling toch betere input kan zijn voor het voorraadbeheer. Hoewel voorraadbeheer buiten de scope van dit onderzoek valt, is vanwege extra tijd, besloten om (op basis van de literatuur) toch een klein bruggetje te slaan. Na dit onderzoek zal een andere afstudeerder verder in gaan op het voorraadbeheer en de interactie tussen *demand forecasting* en *stock control*.

Interactie voorspelling en voorraadbeheer

Onze vermoedens dat de voorspelnauwkeurigheid niet alles zegt over de prestaties van het voorraadbeheer wordt door de literatuur bevestigd. Zo komt naar voren dat er geen rechtstreekse relatie is tussen de prestaties van voorspellingen (aan de hand relevante prestatie-indicatoren) en de effectiviteit van het voorraadbeheer (Syntetos, Nikolopoulos, & Boylan, 2010). Daarnaast is het belangrijk om de interactie tussen voorspelling en voorraadbeheer te begrijpen omdat de prestaties van het voorraadbeheer wordt bepaald door beide componenten (Ali, Boylan, & Syntetos, 2012), (Strijbosch, Syntetos, Boylan, & Janssen, 2011) & (Tiacchi & Saetta, 2009). In veel van de onderzochte literatuur wordt de voorspelling als doel op zich beschouwd, terwijl vanuit een praktisch oogpunt de voorspelling juist als middel wordt gebruikt om het uiteindelijke doel te bereiken (optimaal voorraadbeheer). Om de prestaties van voorraadbeheer te bepalen wordt vaak gekeken naar voorraadkosten en/of behaalde service levels, dit zijn wezenlijk andere prestatie-indicatoren dan gebruikt bij voorspelaccuratie. Voor een verlenging van dit onderzoek is het daarom erg interessant om te kijken welke variabelen van invloed zijn op de prestaties van het voorraadbeheer. Zo stelt (Beutel & Minner, 2012) bijvoorbeeld dat de hoogte van veiligheidsvoorraad om een bepaalde service level te bewerkstelligen afhangt van de vraagonzekerheid en de bijbehorende voorspelfouten. Er wordt ook een parallel getrokken tussen veiligheidsvoorraad en de variatiecoëfficiënt. Zo zou er bijvoorbeeld een relatie kunnen zijn tussen een hoge variatiecoëfficiënt (erratic en lumpy artikelen) en de hoogte van de veiligheidsvoorraad om aan een bepaald servicelevel te voldoen. Ook de mate van data-aggregatie (bv. maand of kwartaal) kan van invloed zijn op service levels en voorraadkosten (Babai, Ali, & Nikolopoulos, 2012). Daarnaast kunnen artikelen met zeer veel nul-vraag (zeer hoge ADI) tot *inventory obsolescence* leiden met het gebruik van Croston of SBA (Teunter, Syntetos, & Babi, 2011). Wellicht kan er ook gebruikt gemaakt worden van de Periods in Stock omdat dit een maat is die zich richt op de consequenties van een foutieve voorspelling voor het voorraadbeheer. Een idee van (Wallström & Segerstedt, 2010) is om een 'mean forecasted stock' te gebruiken om de prestaties van de voorspelling te monitoren. Dit zou als volgt gedefinieerd zijn:

$$\text{Mean forecasted stock} = \frac{PIS_T}{T}$$

Het is duidelijk dat er genoeg aanknopingspunten zijn voor verder onderzoek. De relatie begrijpen tussen *demand forecasting* en *stock control* is van essentieel belang om een goed voorraadbeheer op te bouwen (Altay, Litteral, & Rudisill, 2012). De literatuur over dit onderwerp (zover beschouwd) is inmiddels doorgegeven aan de volgende afstudeerder.

Invloed op voorraadbeheer

Om een klein luikje naar voorraadbeheer te openen en om ons vermoeden dat de voorspelmethodeken een betere input zijn voor het voorraadbeheer te bevestigen, is hier een kleine test ontworpen. Op basis van een simpele bestelpolitiek is een voorraadbeheer opgesteld voor enkele artikelen. Deze voorbeelden geven geen realistisch beeld van het daadwerkelijke voorraadbeheer en de voorraadhoogtes, maar geven een beeld tussen het verschil in voorspelaccuratie tussen de voorspelmethodeken en de huidige voorspelling en de invloed daarvan op het voorraadbeheer.

We hebben beredeneerd dat de niet significant betere voorspelmethodeken toch van meer toegevoegde waarde kunnen zijn voor de voorraadaansturing. We nemen daarom de proef op de som door op basis van de voorspelmethodeken en op basis de huidige voorspelling de voorraadhoogtes te bepalen. We zijn hierbij uitgegaan van een simpele politiek die er voor zorgt dat er genoeg voorraad is om aan de vraag binnen de levertijd van het artikel te voldoen. Daarbij fungeert een bestelpunt (dat reeds door Stork wordt gehanteerd) als een soort veiligheidsvoorraad: als de voorraadhoogte onder dit punt dreigt te raken binnen de levertijd, wordt er ook een bestelling gedaan om dit aan te vullen tot het bestelpunt. Daarnaast wordt er gebruik gemaakt van de huidige minimale bestelgroottes. Er wordt dus allereerst gekeken wat de voorspelde vraag is binnen de levertijd (op basis van de methodeken en op basis van de huidige voorspelling). Als er niet genoeg voorraad is, wordt er een bestelling geplaatst zodat er wel aan de voorspelde vraag voldaan kan worden, rekening houdend met het (minimale) bestelpunt.

Om de toegevoegde waarde van de voorspelmethodeken aan te tonen, is gekozen om enkele artikelen te gebruiken waar de voorspelmethodek slechter presteert dan de huidige voorspelling (dus een hogere A-MAPE). Vervolgens werd gekeken naar wat de gemiddelde voorraadhoogte was gebaseerd enerzijds op de voorspelmethodek en anderzijds gebaseerd op de huidige voorspelling. Hieruit bleek dat de voorspelmethodeken een gemiddelde lagere voorraad geven, terwijl de A-MAPE juist hoger is. Hiermee is ons vermoeden in eerste instantie bevestigd; de voorspelmethodeken geven betere input voor het voorraadbeheer dan de huidige voorspelling. Dit zijn echter slechts enkele voorbeelden, het totale voorraadbeheer dient nog verder onderzocht te worden. In onderstaande tabel staan drie voorbeelden waar de A-MAPE van de voorspelmethodek hoger is, maar dat dit toch zorgt voor een gemiddeld lagere voorraad⁶. De dikgedrukte waardes geven de beste waardes aan. In bijlage 9 zijn schematische weergaves te vinden van de drie voorbeelden.

Tabel 22 - Invloed verschil tussen huidige voorspelling en voorspelmethodek op voorraadbeheer

	Methodiek	A-MAPE	Gem. Voorraad	Voorraadreductie
Artikel 1	SBA (erratic)	71%	€10.062	19%
	Huidige voorspelling	79%	€12.614	
Artikel 2	SBA (lumpy)	118%	€244	22%
	Huidige voorspelling	105%	€314	
Artikel 3	SBA (erratic)	85%	€2.249	28%
	Huidige voorspelling	65%	€3.125	

⁶ Ook met de lagere voorraad kon aan evenveel vraag voldaan worden.

Hoofdstuk 10 De relatie tussen klantinformatie en voorspelbaarheid

Dit hoofdstuk gaat kort in op de relatie tussen informatie van de klant en de verbetering op het gebied van voorspelling.

Installed base forecasting

De kwantitatieve aanpak zoals gepresenteerd in dit onderzoek voor de vraagvoorspelling van spare parts kent natuurlijk haar grenzen. De historische analyse hoeft geen waarheidsgetrouw beeld te geven aangezien andere (kwalitatieve) factoren zoals marktbewegingen, een groeiende installed base of inspanningen van service personeel ook invloed kunnen zijn op de vraag naar spare parts. Ook uit een literatuuronderzoek blijkt dat het gebruiken van dergelijke informatie aan te bevelen is (Haberletiner, Meyr, & Taudes, 2010). Er is echter weinig literatuur beschikbaar over de rol van klantinformatie op de logistiek voor spare parts. Een recentelijk onderzoek brengt hier echter verandering in door te spreken van *installed base forecasting* (Dekker, Pince, Zuidwijk, & Jalil, 2012). Deze manier van voorspellen legt relatie met informatie van de klanten en de status van de installed base met de te verwachten vraag naar spare parts. Relevante informatie zou bijvoorbeeld zijn: machine type (en daarmee een lijst van onderdelen), installatie datum, type van service contract en staat van de machine (Dekker, Pince, Zuidwijk, & Jalil, 2012). Er wordt bijvoorbeeld een relatie gelegd tussen de *product life cycle* (PLC) en spare parts voorspelling. Aan het begin van de levensfase van een nieuwe machine (waar dus geen historische data beschikbaar is) kan voorspeld worden op basis van *failure rates*. Deze rates zouden geschat kunnen worden aan de hand van eerdere generaties of de praktijkervaring die in de loop van de jaren ontstaan is. Gedurende de levensfase komt geleidelijk aan historische data beschikbaar en daarmee kan de voorspelling worden aangevuld. Aan het eind van de levensfase kan informatie over de installed base (zoals de staat van de machine) een indicator zijn voor de vraag naar spare parts tot en met het einde van de levensduur van de machine. De relatie tussen de PLC en spare parts is een onderzoek op zich.

Het gebruik van *installed base forecasting* kan in theorie de voorspelaccuratie stukken verbeteren omdat er extra wijsheid in de, op historie gebaseerde, voorspelling wordt gebracht in termen van de echte status van de machine. Er zou bijvoorbeeld een correlatie gevonden kunnen worden tussen de vraag naar onderdelen en het aantal reparaties/onderhoudsbeurten (Romeijnders, Teunter, & van Jaarsveld, 2012). Een voorbeeld hiervan bij Stork is de 100SIP-kit. Dit is een kit met te vervangen onderdelen als een machine 100 sterilisatiebeurten heeft gedaan. Het gebruiken van installed base informatie (zoals de SIP-tellers die Stork op afstand kan uitlezen) heeft grote toegevoegde waarde voor de voorspelling. De aard van de vraag(karakteristieken) kan daarmee beter getypeerd worden dan met alleen een historische analyse (Wang & Syntetos, 2011). Daarnaast zou het bijhouden van uitgebreide service rapporten over de staat van de machine, (nog te) vervangen onderdelen, mate van onderhoud, operationele omstandigheden e.d. van toegevoegde waarde kunnen zijn om met deze informatie relaties te leggen met de vraag naar (diverse categorieën) spare parts. De relatie tussen installed base informatie en de vraag naar spare parts dient nog nader onderzocht te worden.

Toch kan het een uitdaging zijn om deze informatie van de klant te krijgen en betekenisvol te maken. Met een grote (en oude) installed base kan het veel tijd en geld kosten om al deze informatie beschikbaar te krijgen. Desondanks heeft deze informatie toegevoegde waarde om betere service te verlenen naar de klant. De interactie tussen de klanten en Stork is belangrijk voor beide partijen en zou bijvoorbeeld kunnen worden vastgelegd in een Service Level Agreement (SLA).

Hoofdstuk 11 Conclusies en aanbevelingen

In dit verslag is onderzoek gedaan naar de vraagvoorspelling van spare parts aan de hand van vijf onderzoeksvragen. Per onderzoeksvraag worden de conclusies behandeld, daarna volgen enkele niet-onderzoeksvraag specifieke conclusies en de aanbevelingen.

Onderzoeksvraag 1

Hoe is de huidige wijze van vraagvoorspelling en wat is de kwaliteit hier van?

De huidige manier van voorspellen wordt gedaan aan de hand van een vierjaarlijks gemiddelde. De vraag per jaar wordt geacht uniform verdeeld te zijn over het jaar. De gemiddelde voorspelfout (A-MAPE) bleek 96,64% te zijn, de Number of Shortages (NOSp) 39,80% en uit de lineaire regressie bleek dat de cumulatieve voorspelling sneller stijgt dan de cumulatieve vraag.

Onderzoeksvraag 2

Hoe kunnen spare parts op een slimme manier worden gecategoriseerd ten behoeve van vraagvoorspelling?

In dit onderzoek is gekozen om de classificatie te baseren op twee variabelen: 1) de gemiddelde tussentijd en 2) de gekwadrateerde variatiecoëfficiënt van de vraaghoeveelheid.

Uit deze classificatie komen vier groepen: 1) smooth, artikelen die een lage tussentijd hebben en een lage variatiecoëfficiënt, 2) erratic, artikelen met een lage tussentijd maar met een hoge variatiecoëfficiënt, 3) intermittent, artikelen met een hoge tussentijd en een lage variatiecoëfficiënt en 4) lumpy, artikelen met zowel een hoge tussentijd alsmede een hoge variatiecoëfficiënt.

Onderzoeksvraag 3

Welke voorspellingsmethodieken zijn geschikt voor de diverse categorieën spare parts?

Er is gekozen om in de categorie smooth de methode van Croston toe te passen en in de andere drie categorieën de Syntetos & Boylan Approximation.

Onderzoeksvraag 4

Hoe kunnen geselecteerde voorspellingsmethodieken getoetst worden?

De prestaties van de methodieken zijn gemeten aan de hand van verschillende prestatie-indicatoren, namelijk de Mean Squared Error (MSE), de Adjusted Mean Absolute Percentage Error (A-MAPE), de Cumulated Forecast Error (CFE) de Number of Shortages (NOSp) en de Periods in Stock (PIS). Er is gekozen voor de A-MAPE in plaats van de MAPE omdat de eerstgenoemde om kan gaan met nul-vraag (iets dat regelmatig voorkomt bij spare parts).

Onderzoeksvraag 5

Hoe kunnen de geselecteerde voorspellingsmethodieken worden onderhouden?

Door een softwarematige oplossing die regelmatig (maandelijks) vraagpatronen analyseert en (grote) veranderingen op een slimme manier in kaart weet te brengen zou zowel een praktische als een tijdefficiënte manier zijn om op een dynamischere manier de voorspelling op artikelniveau up-to-date te houden.

Het bleek dat de smooth categorie het beste te voorspellen is (met een gemiddelde A-MAPE van 55,93%), vervolgens erratic (71,93%), daarna intermittent (103,80%) en tenslotte lumpy (127,18%).

Uit een vergelijking tussen de huidige manier van voorspellen en de voorspelmethodeken bleek echter dat er geen verbeterde voorspelaccuratie te vinden was (in termen van MSE, A-MAPE en PIS). De implicaties van de voorspelmethodeken voor het voorraadbeheer zijn echter nog niet onderzocht, maar de verwachting is dat de voorspelmethodeken (vanwege hun meer dynamischere karakter) een betere input zijn voor de voorraadbeheerder dan de huidige uniforme benadering. Er kan pas met zekerheid worden gezegd hoe of de voorspelmethodeken toegevoegde waarde hebben als de interactie tussen voorspelling en voorraadbeheer nader is onderzocht.

Daarnaast is ook laten zien dat door de voorspelling aan te passen op basis van niet-historische gegevens er een enorme winst behaald kan worden. Dit kan bijvoorbeeld door van te voren bekende vraag als order in te leggen. Het gebruik en de consequenties van andere (kwalitatieve) informatie voor de voorspelling dient nog nader onderzocht te worden.

Aanbevelingen

- Artikelen dienen geclassificeerd te worden ten behoeve van vraagvoorspelling aan de hand van tussentijdintervallen (ADI) en de variatie in bestelhoeveelheid (CV^2). De scheidingswaarde ligt bij de ADI op 1,32 en bij de CV^2 op 0,49. Er moet ook gekozen worden of een artikel wel of niet voorspeld gaat worden (bijvoorbeeld aan de hand van minimum aantal vraagmomenten).
- Voor de categorie smooth is de Croston methode het meest geschikt, voor de andere categorieën zou gekozen moeten worden voor de Syntetos & Boylan Approximation.
- Er dient maandelijks voorspeld te worden (verder vooruit al naar gelang de levertijd).
- Er dient zoveel mogelijk deterministische informatie in het voorspelproces ingebracht te worden, bijvoorbeeld door het inleggen van orders. Nauwe samenwerking met de klant en betere benutting van beschikbare data van de installed base kan hier van zeer toegevoegde waarde zijn. De aanpassing van de voorspelling dient echter met wijsheid te gebeuren. Degene moet zijn keuzes kunnen verantwoorden zodat de voorspelling niet uit de pas gaat lopen.
- Er dient nader onderzocht te worden hoe het voorraadbeheer optimaal af te stemmen is op de voorspelling.
- Om goed grip te krijgen op de voorspelling wordt ten zeerste aanbevolen om dit softwarematig te doen. Bijvoorbeeld door gebruik te maken van het SAP ERP systeem of een uitbreidingspakket zoals Slim4. Een eventueel uitbreidingspakket dient naadloos samen te werken met het huidige SAP ERP systeem.
- Het softwarepakket dient zo ingericht te zijn dat de meeste artikelen automatisch voorspeld worden en dat met behulp van enkele uitzonderingsregels gemakkelijk overzichten naar voren te halen zijn met nader te bestuderen artikelen. Daardoor komt de meeste aandacht te liggen bij de artikelen die ook daadwerkelijk aandacht verdienen.

Bibliografie

- Ali, M. M., Boylan, J. E., & Syntetos, A. A. (2012). Forecast errors and inventory performance under forecast information. *International Journal of Forecasting*, 28, 830-841.
- Altay, N., Litteral, L. A., & Rudisill, F. (2012). Effects of correlation on intermittent demand forecasting and stock control. *International Journal of Production Economics*, 135, 275-283.
- Amin-Naseri, M. R., & Tabar, B. R. (2008). Neural Network Approach to Lumpy Demand Forecasting For Spare Parts in Process Industries. *Proceedings of the International Conference on Computer and Communication Engineering*, (pp. 1378-1382). Kuala Lumpur, Malaysia.
- Archibald, B. C., & Koehler, A. B. (2003). Normalization of seasonal factors in Winter's methods. *International Journal of Forecasting*, 19, 143-148.
- Babai, Z. M., Ali, M. M., & Nikolopoulos, K. (2012). Impact of temporal aggregation on stock control performance of intermittent demand estimators: Empirical analysis. *Omega*, 40, 713-721.
- Bacchetti, A., & Saccani, N. (2012). Spare parts classification and demand forecasting for stock control: Investigating the gap between research and practice. *Omega*, 40, 722-737.
- Beutel, A.-L., & Minner, S. (2012). Safety stock planning under casual demand forecasting. *International Journal of Production Economics*, 140, 637-645.
- Boylan, J. E., & Syntetos, A. A. (2010). Spare parts management: a review of forecasting research and extensions. *IMA Journal of Management Mathematics*, 21(3), 227-237.
- Boylan, J. E., Syntetos, A. A., & Karakostas, G. C. (2008). Classification for Forecasting and Stock Control: A Case Study. *The Journal of the Operational Research Society*, 59(4), 473-481.
- BuHamra, S., Smaoui, N., & Gabr, M. (2003). The Box-Jenkins analysis and neural networks: prediction and time series modeling. *Applied Mathematical Modelling*, 27, 805-815.
- Callegaro, A. (2010). *Forecasting Methods for Spare Parts Demand*. Opgeroepen op November 13, 2012, van Padoba Digital University Archive: <http://tesi.cab.unipd.it/25014/1/TesiCallegaro580457.pdf>
- Cavalieri, S., Garetti, B., Macchi, M., & Pinto, R. (2008). A decision-making framework for managing maintenance spare parts. *Production Planning & Control*, 19(4), 379-396.
- Chen, F.-L., & Chen, Y.-C. (2009). An investigation of Forecasting Critical Spare Parts Requirement. *Word Congress on Computer Science and Information Engineering*, (pp. 225-230). Los Angeles, CA.
- Chen, F.-L., Chen, Y.-C., & Kuo, J.-Y. (2010). Applying Moving back-propagation neural network and Moving fuzzy-neuron network to predict the requirement of critical spare parts. *Expert Systems with Applications*, 37, 6695-6704.
- Cohen, M. A., Agrawal, N., & Agrawal, V. (2006, Mei). Winning in the Aftermarket. *Harvard Business Review*, 84(5), 129-138.

- Croston, J. D. (1972). Forecasting and stock control for intermittent demands. *Operational Research Quarterly*, 23(3), 289-303.
- Dekker, R., Pince, C., Zuidwijk, R., & Jalil, M. N. (2012). On the use of installed base information for spare parts logistics: A review of ideas and industry practice. *International Journal of Production Economics*, doi:10.1016/j.ijpe.2011.11.025.
- Driessen, M. A., Arts, J. J., van Houtum, G. J., Rustenburg, W. D., & Huisman, B. (2012, April 11). Maintenance spare parts planning and control: A framework for control and agenda for future research. pp. 1-31.
- Ghobbar, A. A., & Friend, C. H. (2002). Sources of intermittent demand for aircraft spare parts within airline operations. *Journal of Air Transport Management*, 8, 221-231.
- Ghobbar, A. A., & Friend, C. H. (2003). Evaluation of forecasting methods for intermittent parts demand in the field of aviation: a predictive model. *Computers & Operations Research*, 2097-2114.
- Haberletiner, H., Meyr, H., & Taudes, A. (2010). Implementation of a demand planning system using advance order information. *International Journal of Production Economics*, 128, 518-526.
- Heinecke, G., Syntetos, A. A., & Wang, W. (2011, November 23). Forecasting-based SKU classification. *International Journal of Production Economics*.
- Hu, H.-Y., & Li, Z.-C. (2006). Collocation methods for Poisson's equation. *Computer Methods in Applied Mechanics and Engineering*, 195, 4139-4160.
- Kostenko, A. V., & Hyndman, R. J. (2006, Februari 14). *A note on the categorization of demand patterns*. Opgeroepen op November 13, 2012, van Rob J Hyndman: <http://www.robjhyndman.com/papers/idcat.pdf>
- Romeijnders, W., Teunter, R., & van Jaarsveld, W. (2012). A two-step method for forecasting spare parts demand using information on component repairs. *European Journal of Operational Research*, 220, 386-393.
- Silver, E. A., Pyke, D. F., & Peterson, R. (1998). *Inventory Management and Production Planning and Scheduling* (3rd ed.). New York: John Wiley & Sons.
- Strijbosch, L. W., Syntetos, A. A., Boylan, J. E., & Janssen, E. (2011). On the interaction between forecasting and stock control: The case of non-stationary demand. *International Journal of Production Economics*, 133, 470-480.
- Syntetos, A. A., & Boylan, J. E. (2001). On the bias of intermittent demand estimates. *International Journal of Production Economics*, 71(1-3), 457-466.
- Syntetos, A. A., & Boylan, J. E. (2005). The accuracy of intermittent demand estimates. *International Journal of Forecasting*, 21, 303-314.
- Syntetos, A. A., Boylan, J. E., & Croston, J. D. (2005). On the categorization of Demand Patterns. *The journal of the Operational Research Society*, 56(5), 495-503.

- Syntetos, A. A., Nikolopoulos, K., & Boylan, J. E. (2010). Judging the judges through accuracy-implication metric: The case of inventory forecasting. *International Journal of Forecasting*, 26, 134-143.
- Syntetos, A. A., Nikolopoulos, K., Boylan, J. E., Fildes, R., & Goodwin, P. (2009). The effect of integrating management judgement into intermittent demand forecasts. *International Journal of Production Economics*, 118, 72-81.
- Teunter, R. H., Syntetos, A. A., & Babai, M. Z. (2010). Determining order-up-to levels under periodic review for compound binomial (intermittent) demand. *European Journal of Operational Research*, 203, 619-624.
- Teunter, R. H., Syntetos, A. A., & Babi, M. Z. (2011). Intermittent demand: Linking forecasting to inventory obsolescence. *European Journal Of Operational Research*, 214, 606-615.
- Tiacci, L., & Saetta, S. (2009). An approach to evaluate the impact of interaction between demand forecasting methods and stock control policy on the inventory system performances. *International Journal of Production Economics*, 118, 63-71.
- Tseng, F.-M., Yu, H.-C., & Tzeng, G.-H. (2002). Combining neural network model with seasonal time series ARIMA model. *Technological Forecasting & Social Change*, 69, 71-87.
- van Kampen, T. J., Akkerman, R., & van Donk, D. (2012). SKU classification: A literature review and conceptual framework. *International Journal of Operations and Production Management*, 32(7), 850-876.
- Wallström, P., & Segerstedt, A. (2010). Evaluation of forecasting error measurements and techniques for intermittent demand. *International Journal of Production Economics*, 128, 625-636.
- Wang, W., & Syntetos, A. A. (2011). Spare parts demand: Linking forecasting to equipment maintenance. *Transportation Research Part E: Logistics and Transportation Review*, 47(6), 1194-1209.
- Willemain, T. R., Smart, C. N., & Schwarz, H. F. (2004). A new approach to forecasting intermittent demand for service parts inventories. *International Journal of Forecasting*, 20, 375-387.
- Williams, T. M. (1984). Stock Control with Sporadic and Slow-Moving Demand. *The Journal of the Operational Research Society*, 35(10), 939-948.
- Yar, M., & Chatfield, C. (1990). Prediction intervals for the Holt-Winters forecasting procedure. *International Journal of Forecasting*, 6, 127-137.

Bijlagen

Bijlage 1 – Prestatie-indicatoren	43
Bijlage 2 – Overzicht voorspelmethodeken.....	45
Bijlage 3 – Uitwerking bias-gecorrigeerde factor SBA.....	48
Bijlage 4 – Data analyse (1)	49
Bijlage 5 – Data analyse (2)	50
Bijlage 6 – Uitwerking Syntetos-Boylan Approximation	51
Bijlage 7 – Grafische weergave per categorie	53
Bijlage 8 – Overzicht resultaten van alle artikelen	61
Bijlage 9 – Voorbeelden invloed op voorraadbeheer	66
Bijlage 10 – Lijst met figuren en tabellen	67

Bijlage 1 – Prestatie-indicatoren

Naam	Afkorting	Definitie	Opmerkingen
Mean Error	ME	$ME = \frac{1}{n} \sum_{t=1}^N (F_t - X_t)$	ME = 0 betekent niet perfecte voorspelling, plus en min kunnen elkaar wegstrepen.
Mean Absolute Error	MAE	$MAE = \frac{1}{n} \sum_{t=1}^N (F_t - X_t)$	Behandelt over-voorspelling en under-voorspelling gelijk.
Mean Squared Error	MSE	$MSE = \frac{1}{n} \sum_{t=1}^N (F_t - X_t)^2$	Schaalafhankelijk. Gevoelig voor outliers en $F_t < 1$, vanwege het kwadraat. Geeft meer gewicht aan grote fouten.
Mean Absolute Percentage Error	MAPE	$MAPE = \frac{1}{n} \sum_{t=1}^n \left \frac{F_t - X_t}{X_t} \right $	Niet gedefinieerd voor $X_t = 0$, enorm scheve verdeling voor $X_t \approx 0$.
Scaled Mean Absolute Percentage Error	S-MAPE	$S - MAPE = \frac{1}{n} \sum_{t=1}^n \frac{ F_t - X_t }{(X_t + F_t)/2}$ $S - MAPE (100\%) = \frac{1}{n} \sum_{t=1}^n \frac{ F_t - X_t }{(X_t + F_t)}$	Niet schaalafhankelijk, maar behandelt over-voorspelling en under-voorspelling niet gelijk. Geeft een waarde tussen -200% en +200%. De 2 ^{de} formule tussen -100% en +100%
Adjusted Mean Absolute Percentage Error	A-MAPE	$A - MAPE = \frac{\frac{\sum_{t=1}^n F_t - X_t }{N}}{\frac{\sum_{t=1}^n X_t}{N}}$	Aangepaste variant van MAPE die wel met nul-vraag om kan gaan.
Root Mean Squared Error	RMSE	$RMSE = \sqrt{MSE}$	Voordeel ten opzichte van MSE is dat de schaal van RMSE meer overeenkomt met de data.
Cumulated Forecast Error	CFE	$CFE_t = \sum_{i=1}^t (X_i - F_i) = X_t - F_t + CFE_{t-1}$ $CFE_{max} = \max_{t \in \{1,2,...,T\}} CFE_t$ $CFE_{min} = \min_{t \in \{1,2,...,T\}} CFE_t$	Geeft een beeld van systematische voorspelfout. Streept min en plus tegen elkaar weg; daarom 2 extra maten. Een systematisch positieve voorspelling geeft negatieve CFE-waarden.

Naam ⁷	Afkorting	Definitie	Opmerkingen
Number of Shortages	NOS	$\text{Als } X_t \neq 0 \text{ \& } CFE_t > 0 \text{ dan } NOS \leftarrow NOS + 1$ $NOS_p = \frac{NOS}{N} 100$	Fictief aantal perioden met tekorten (zonder voorraadbeheer). Indicator voor systematische fout. NOS_p is een verhoudings-versie.
Periods in Stock	PIS	$PIS_t = PIS_{t-1} + \sum_{i=1}^t (X_t - F_t) = \sum_{i=1}^t (X_t - F_t)(t + 1 - i)$ $PIS_T = - \sum_{t=1}^T CFE_t$	Geeft een maat voor het verschil tussen voorspelling en werkelijkheid, maar ook hoelang het duurt voordat dit gecorrigeerd is.

⁷ Beide maten zijn een relatief nieuw begrip in de literatuur om voorspelaccuratie te meten. Beide maten werden voor het eerst genoemd in (Wallström & Segerstedt, 2010)

Bijlage 2 – Overzicht voorspelmethodeken

Methodiek	Afk.	Variabelen	Omschrijving	Voordelen	Nadelen	Referenties
Single Exponential Smoothing	SES	- Historische data - Smoothing constante	Combinatie tussen lineair patroon historische data en de meest recente data	- Makkelijk te berekenen - Geschikt voor smooth demand	- Ongeschikt voor perioden met nul-vraag en hoge variatie	(Syntetos & Boylan, 2005) (Ghobbar & Friend, 2002) (Willemain, Smart, & Schwarz, 2004)
Moving Average	MA	- Historische data - Aantal te beschouwen perioden	Gemiddelde van de aantal te beschouwen perioden	- Makkelijk te berekenen - Geschikt voor smooth demand	- Behandelt alle historische waardes gelijk	(Chen & Chen, 2009)
Weighted Moving Average	WMA	- Historische data - Aantal te beschouwen perioden	Gemiddelde van de aantal te beschouwen perioden met afnemende gewichten	- Makkelijk te berekenen - Meer gewicht op laatste periode	- Alleen geschikt als vraag niet lumpy is	(Ghobbar & Friend, 2003)
Croston's Method	CR	- Historische data - Interval tussen heden en vorige niet-nul vraag - Smoothing constante	Verbetering van SES, kijkt ook naar interval van nul-vraag	- Toepasbaar op vraag met (veel) nul-perioden	- Bevat theoretische fout	(Croston, 1972)
Syntetos-Boylan Approximation	SBA	- Historische data - Interval tussen heden en vorige niet-nul vraag - Smoothing constante	Verbetering van CR zonder systematische fout	- Toepasbaar op vraag met (veel) nul-perioden - Geen theoretische fout	- Eén van de verbeteringen van CR	(Syntetos & Boylan, 2005) (Syntetos & Boylan, 2001) (Wallström & Segerstedt, 2010) (Amin-Naseri & Tabar, 2008)

Methodiek	Afk.	Variabelen	Omschrijving	Voordelen	Nadelen	Referenties
Additive Winter	AW	- Historische data - Smoothing constante - Trend constante - Periodieke constant - Breedte van periodiek	Verbetering van SES die een trend toevoegd	- Houdt rekening met een trend (op een toevoegende manier)	- Spare parts zijn over het algemeen niet trend-gevoelig	(Archibald & Koehler, 2003) (Ghobbar & Friend, 2003) (Yar & Chatfield, 1990)
Multiplicative Winter	MW	- Historische data - Smoothing constante - Trend constante - Periodieke constant - Breedte van periodiek	Verbetering van SES die een trend vermenigvuldigd	- Houdt rekening met een trend (op een vermenigvuldigen de manier)	- Spare parts zijn over het algemeen niet trend-gevoelig	(Archibald & Koehler, 2003) (Ghobbar & Friend, 2003)
Poisson method	PM	- Historische data - Intervaltijd - Waarde van vraag om te voorspellen	Toepassing van de Poisson verdeling	- Statistisch in plaats van deterministisch - Toepasbaar in het geval van sporadische vraag	- Voorspelt geen exacte waarde - Overschat categorieën lumpy en erratic	(Hu & Li, 2006)
Binomial method	BM	- Historische data - Intervaltijd - Level of Service	Toepassing van de binomial verdeling	- Statistisch in plaats van deterministisch	- Geeft geen voorspelling maar aantal om aan service level te voldoen	(Teunter, Syntetos, & Babai, 2010)

Methodiek	Afk.	Variabelen	Omschrijving	Voordelen	Nadelen	Referenties
Bootstrap method	Boot	- Historische data - Aantal te nemen samples - Grote van samples	Neemt een x aantal samples uit de data om zo nieuw data te generen	- Toepasbaar wanneer er weinig historische data beschikbaar is	- Kan enorme fouten bevatten door manipulatie data	(Willemain, Smart, & Schwarz, 2004)
(S)-AR(I)MA Box-Jenkins method	BJ	- Historische data - Waarde p voor AR - Waarde q voor MA - Waarde d voor restverschil - In geval van trend: P, D, Q - AR, MA coëfficiënten	Combinatie van een autoregressive (AR) deel en een moving average (MA) deel op een iteratieve manier.	- Kan rekening houden met trend - Iteratief zodat de juiste methode gevonden kan worden	- Veel historische data nodig voor een goed werkend model	(Tseng, Yu, & Tzeng, 2002) (BuHamra, Smaoui, & Gabr, 2003)
Grey prediction model	GM	- Historische data	Maakt gebruik van cumulatieve vraag en de kleinste-kwadraten methode	- Goed toepasbaar wanneer er weinig data beschikbaar is - Goed voorspelling op de korte termijn	- Minder goed in langere termijn voorspelling - Wiskundig ingewikkeld	(Chen & Chen, 2009)
Neural Network	NN	- Historische data - Aantal neuronen - Aantal lagen - Leeralgoritme - Maat voor error	Complex netwerk die met behulp van algoritmes input en output met elkaar koppelt en zo zichzelf traint en slimmer wordt.	- Leert automatisch alle variabelen met elkaar te koppelen	- Zeer complex - Wiskundig ingewikkeld - Veel historische data benodigd - Veel tijd benodigd om het netwerk te definiëren en te trainen	(Chen & Chen, 2010) (Amin-Naseri & Tabar, 2008) (BuHamra, Smaoui, & Gabr, 2003) (Tseng, Yu, & Tzeng, 2002)

Bijlage 3 – Uitwerking bias-gecorrigeerde factor SBA

De volgende berekeningen zijn gebaseerd op (Croston, 1972) en (Syntetos & Boylan, 2005).

Volgens de methode van Croston is de (zuiver) geschatte gemiddelde grootte van de vraag ($Z'_t = E(Z_t) = E(Z'_t) = \mu$) en de (zuiver) geschatte gemiddelde tussentijdinterval ($P'_t = E(P_t) = E(P'_t) = \rho$), zolang geldt dat $X_t > 0$ (waar α voor beide formules gelijk is). Als $X_t = 0$ dan blijven de geschatte parameters hetzelfde als de vorige periode. De voorspelling voor de volgende periode is dan gegeven door $F'_t = \frac{Z'_t}{P'_t}$. In dat geval zou de verwachting van de voorspelling gelijk moeten zijn aan $E(F'_t) = E\left(\frac{Z'_t}{P'_t}\right) = \frac{E(Z'_t)}{E(P'_t)} = \frac{\mu}{\rho}$. Dat zou betekenen dat de methode géén bias heeft.

Syntetos & Boylan (2001) hebben echter aangetoond dat dit wel het geval is, want als wordt verondersteld dat de schatting van de vraag en de schatting van de tussentijdinterval onafhankelijk zijn, dan zou gelden dat:

$$E\left(\frac{Z'_t}{P'_t}\right) = E(Z'_t)E\left(\frac{1}{P'_t}\right), \text{ maar:}$$

$$E\left(\frac{1}{P'_t}\right) \neq \frac{1}{E(P'_t)}$$

En daarom is de methode van Croston biased en dus niet zuiver. Volgens de auteurs komt deze bias niet door de veronderstelling dat de tussentijdintervallen geometrisch verdeeld zijn (zoals aangenomen door Croston). Syntetos & Boylan (2001) hebben een bias-corrigerende factor voorgesteld die bij benadering wél zuiver is. Deze is als volgt opgebouwd:

Laat $x_1 = Z'_t$ en veronderstel dat $E(x_1) = \mu$ en $x_2 = P'_t$ en veronderstel dat $E(x_2) = \rho$.

Met behulp van de stelling van Taylor met de functie $g(x_1, x_2) = \frac{x_1}{x_2}$ en verondersteld dat de verwachting van de vraag en de verwachting van de tussentijdintervallen onafhankelijk zijn dan geldt dat:

$$E(g(x_1, x_2)) = \frac{\mu}{\rho} + \frac{1}{2} \frac{\partial^2 g(\mu, \rho)}{\partial x_2^2} \text{var}(x_2) + \dots + \frac{1}{n} \frac{\partial^n g(\mu, \rho)}{\partial x_2^n}$$

En dus:

$$E(g(x_1, x_2)) = \frac{\mu}{\rho} + \frac{\mu}{\rho^3} + \text{var}(x_2) + \dots$$

Omdat x_2 de exponentieel 'gesmoothde' schatting van de tussentijdinterval is die geometrisch verdeeld is met een variantie $\rho(\rho - 1)$ geldt dat:

$$E\left(\frac{Z'_t}{P'_t}\right) \approx \frac{\mu}{\rho} + \frac{\alpha}{2 - \alpha} \mu \frac{(\rho - 1)}{\rho^2}$$

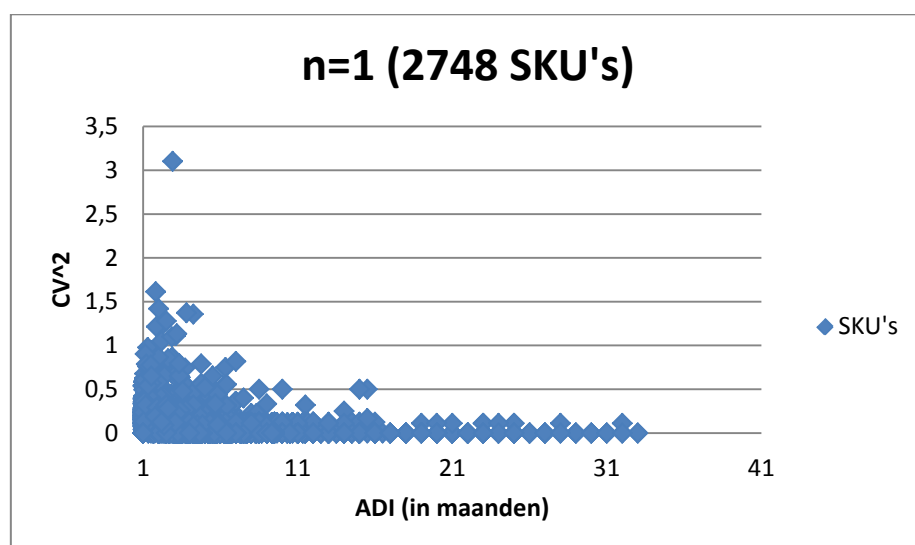
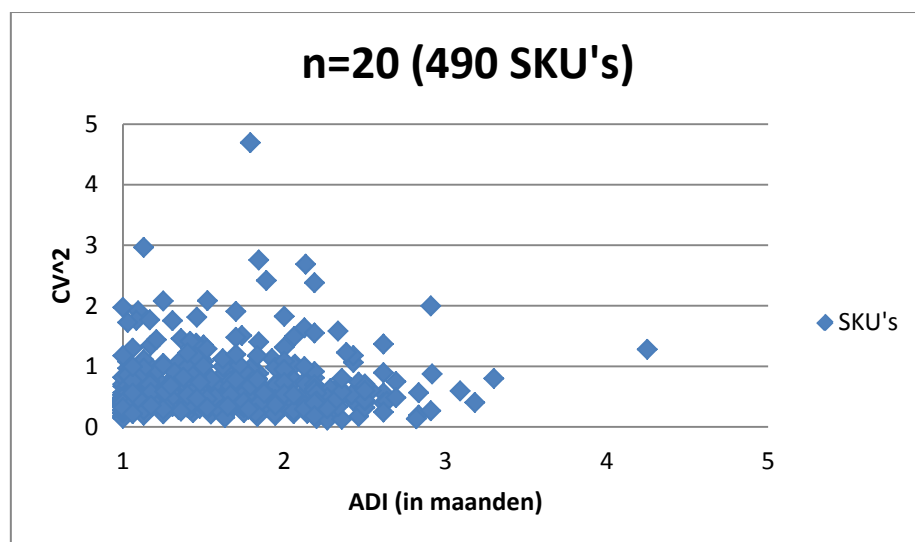
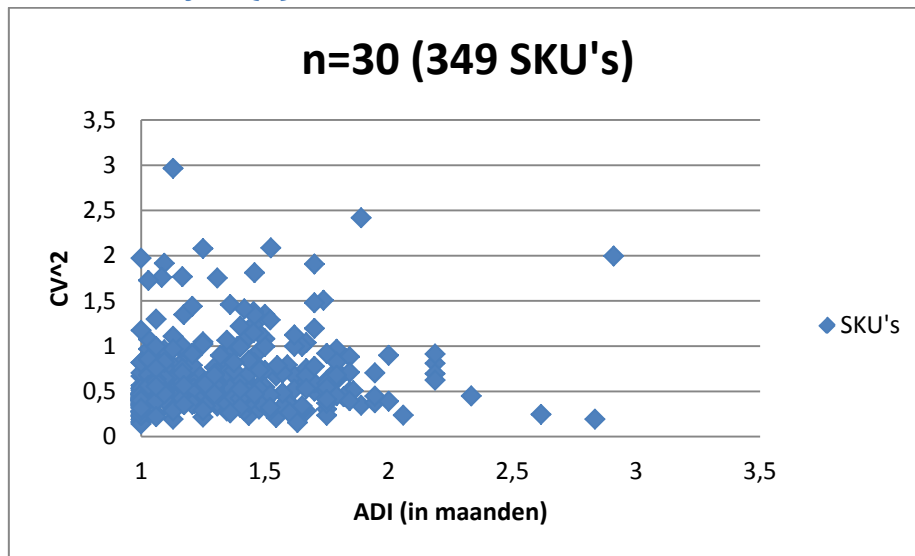
En dus:

$$E\left(\left(1 - \frac{\alpha}{2}\right) \frac{Z'_t}{P'_t}\right) \approx \frac{\mu}{\rho}$$

Bijlage 4 – Data analyse (1)

Door het SAP ERP systeem dat door Stork gehanteerd wordt, was zeer veel data snel beschikbaar en dat vergemakkelijkte de analyse. Na enkele steekproeven op betrouwbaarheid werd allereerst gekeken naar de historie van de afgelopen drie jaar (2011-2009). Dit resulteerde in een bestand dat elke orderregel bevatte voor een spare part samen met meer specifieke informatie als ordernummer, klantnummer, besteldatum en hoeveelheid, inkoopprijs, productgroep, etc. Dit was bedoeld om inzicht te krijgen in hoe het theoretisch kader er in de praktijk zou uitzien onder verschillende voorwaarden. Zo werd gekeken naar verschillende niveaus van minimale aantallen vraagmomenten ($n=30,20,10,1$) over drie jaar en resultaten verder gespecificeerd naar critical/maintenance-ratio en productgroep-ratio. Met behulp van Matlab (R2012b) zijn enkel functies geprogrammeerd die de data automatisch analyseert en classificeert volgens het theoretisch kader. Na dit eerste onderzoek kwamen al enkele belangrijke zaken aan het licht. Zo kon direct worden geconstateerd dat er zeer veel artikelen zijn met zeer weinig vraagmomenten ($n<3$) over drie jaar. Dit was zelfs het merendeel van het aantal artikelen dat verkocht is over deze jaren. Bijlage 5 geeft een beeld van de analyse met verschillende minimum aantal vraagmomenten (n). Dit criterium werd gekozen omdat er artikelen zijn met bijvoorbeeld één vraagmoment over drie jaar en die zijn niet eens de moeite waard om mee te nemen (en dus te voorspellen), deze dienen op een andere manier invulling te krijgen in het voorraadbeheer. Om toch een algeheel beeld te krijgen zijn drie minimum aantal vraagmomenten gepakt; $n=30, 20$ en 1 . De ADI waarde is altijd ≥ 1 (een waarde van 1 duidt op elke maand vraag) en er geldt dat $CV^2 > 0$ (dit om artikelen met géén spreiding; bv. één enkel vraagmoment uit te sluiten).

Bijlage 5 – Data analyse (2)⁸



⁸ n staat voor het minimum aantal vraagmomenten

Bijlage 6 – Uitwerking Syntetos-Boylan Approximation

Hier wordt beschreven hoe de SBA voorspelmethode is geïmplementeerd. Aangezien de methode van SBA hetzelfde werkt als Croston (met uitzondering van de corrigerende factor), geldt deze aanpak ook voor de laatstgenoemde. Alles is geïmplementeerd met behulp van Matlab (R2012b).

Stap 1: definiëren model

$$Z_t = \alpha * X_{t-1} + (1 - \alpha) * Z_{t-1}$$

$$P_t = \alpha * G_{t-1} + (1 - \alpha) * P_{t-1}$$

De variabelen zijn als volgt gedefinieerd:

Z_t = de voorspelling van de omvang van de vraag op tijdstip t

X_t = de werkelijke waarde van vraag op tijdstip t

P_t = de voorspelling van de tijd tussen vraagmomenten op tijdstip t

G_t = de werkelijke waarde van de tijd tussen vraagmomenten op tijdstip t

α = een smoothing constante ($0 \leq \alpha \leq 1$)

De voorspelling voor de vraag op tijdstip $t+1$ wordt dan gegeven door:

$$F_{t+1} = (1 - \frac{\alpha}{2}) \frac{Z_t}{P_t}$$

Stap 2: parameters schatten

Allereerst wordt met behulp van een trainingset de optimale parameters (in dit geval α) van het model geschat. De trainingset bestaat uit de data van 2010 ($t=12$). Om het model te starten wordt verondersteld dat de eerste waarde van de training set perfect voorspeld is zodat $Z_1 = X_1$ en daarnaast wordt verondersteld dat $P_1 = 1$. Zodat $F_1 = Z_1$. In het geval van $X_t = 0$ geldt dat $Z_t = Z_{t-1}$ en $P_t = P_{t-1}$, net zolang totdat $X_t > 0$.

De methodiek is toegepast voor $0 \leq \alpha \leq 1$ met een stapgrootte van 0,05. Bij elke iteratie (aanpassing van α) werd de Sum of Squared Error (SSE)⁹ berekend, gedefinieerd door:

$$SSE = \sum_{t=1}^N (F_t - X_t)^2$$

De waarde van α met de kleinste SSE werd vervolgens gekozen als parameter.

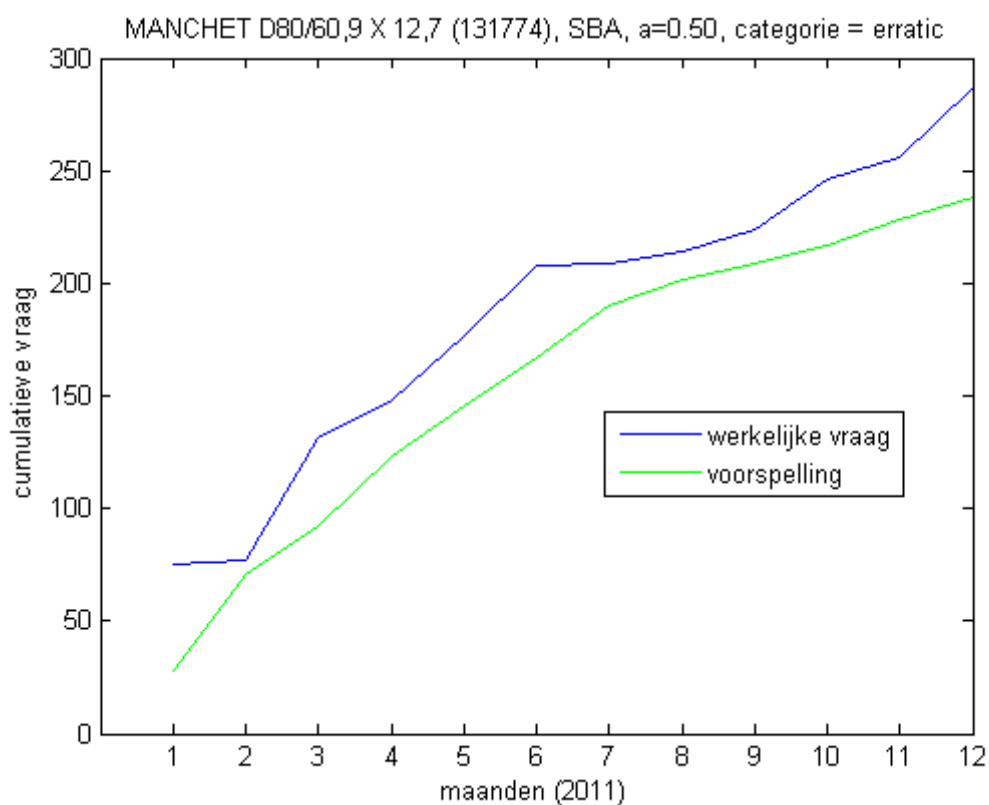
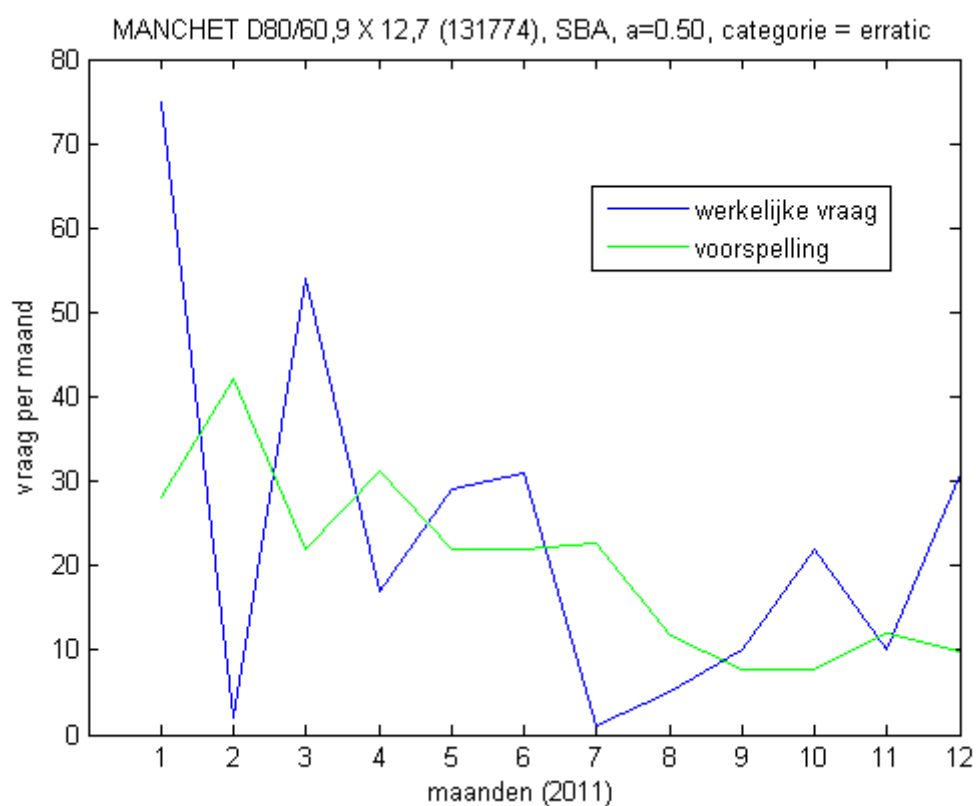
Stap 3: voorspellen

Nadat de meest geschikte parameter is geschat, wordt de testset ingeladen. Deze testset bestaat uit de gegevens van 2011 ($t = 12$). Daarnaast geldt nu dat $F_1 = F_{13-\text{laatste uit testset}}$ en $P_1 = P_{13-\text{laatste uit testset}}$. De methodiek wordt elke maand geüpdatet zodat de vraag van afgelopen maand (die nu bekend is), wordt meegenomen voor de volgende te voorspellen maand.

⁹ Deze maat is schaal-gevoelig, maar de schaal *binnen* de trainingset verandert niet.

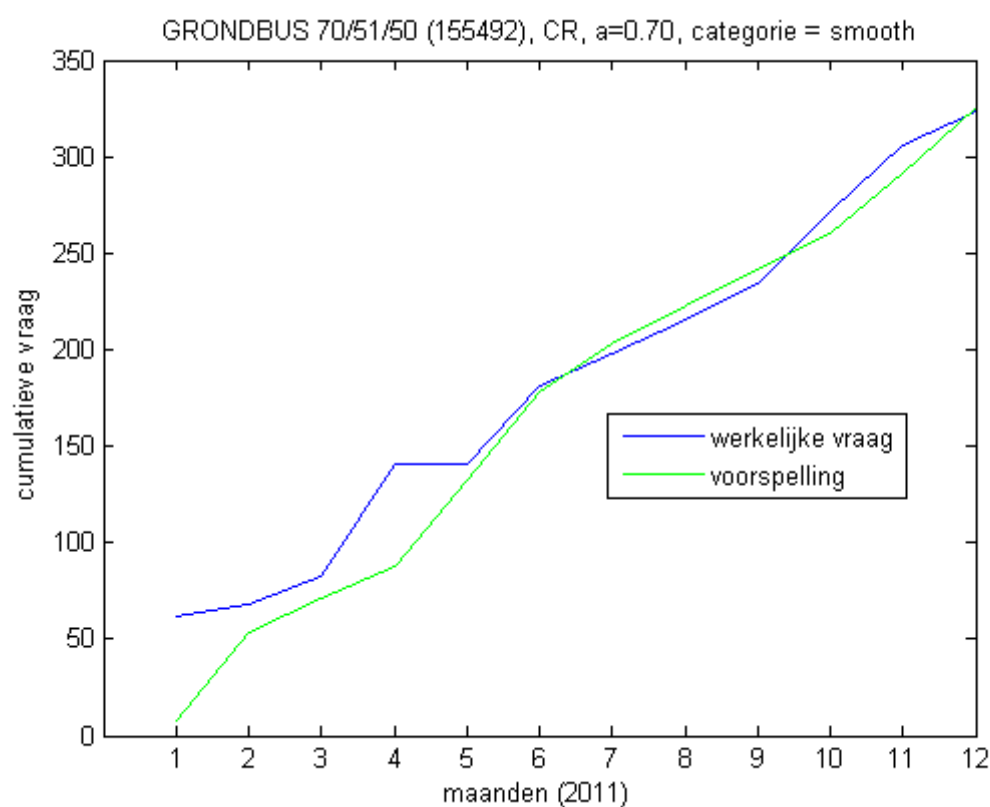
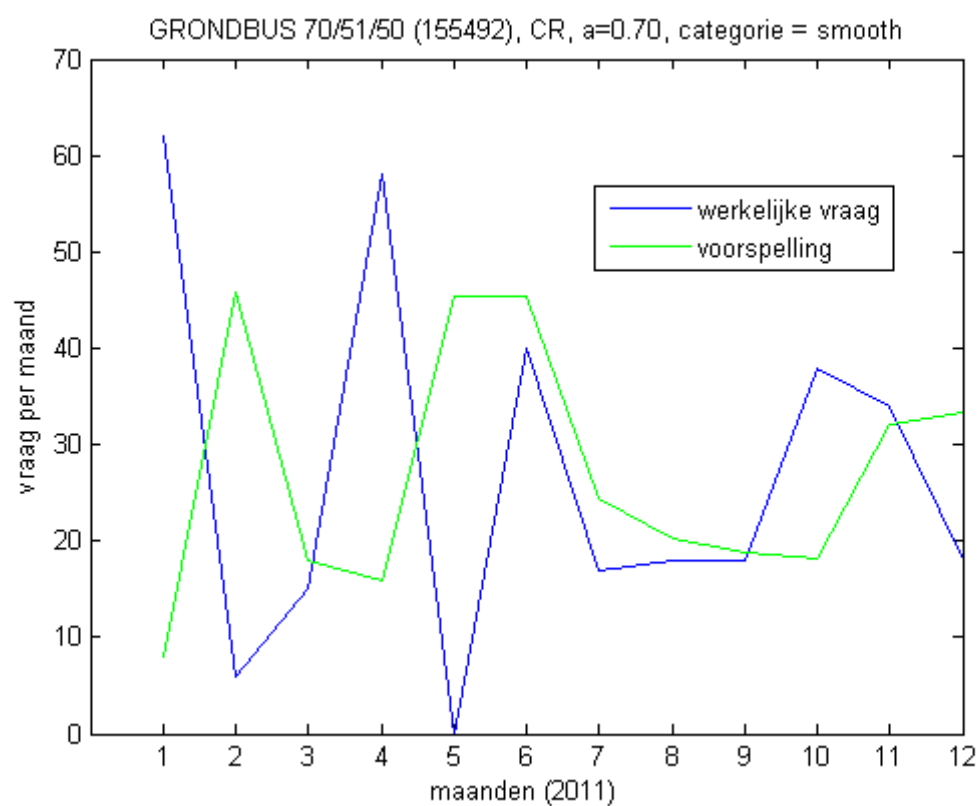
Stap 4: resultaten weergeven

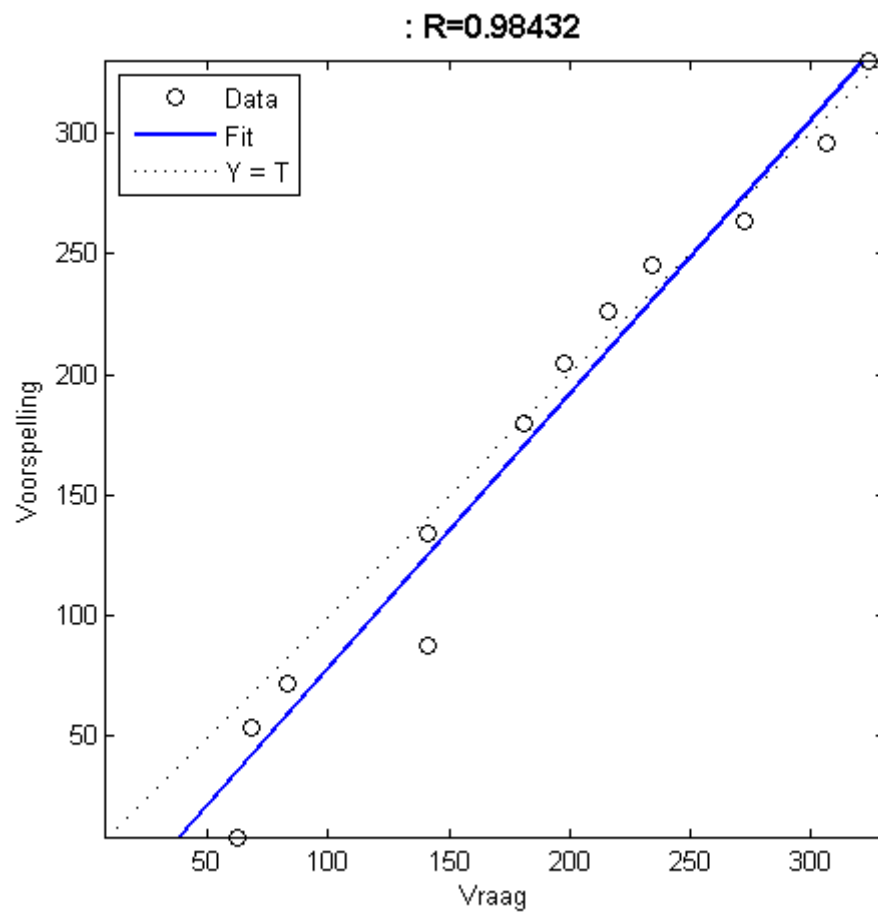
Nadat de methodiek is toegepast wordt een grafische weergave gemaakt van de voorspelling alsmede een cumulatieve voorspelling. Hieronder staat een voorbeeld van zo'n weergave.



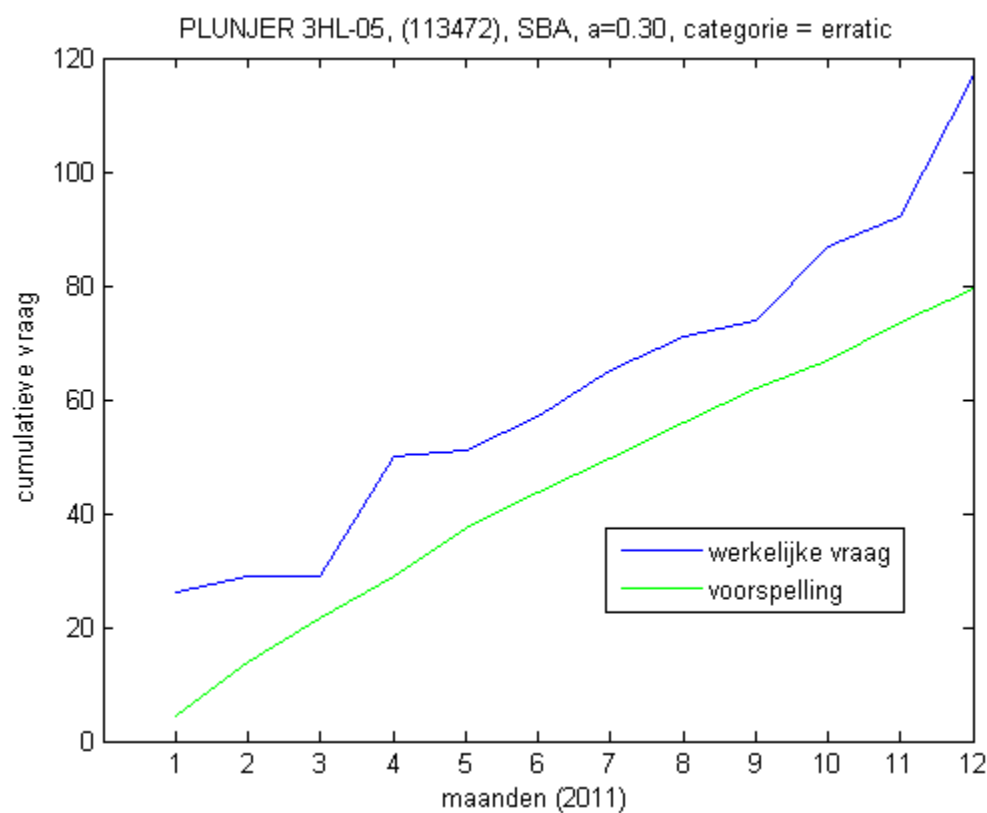
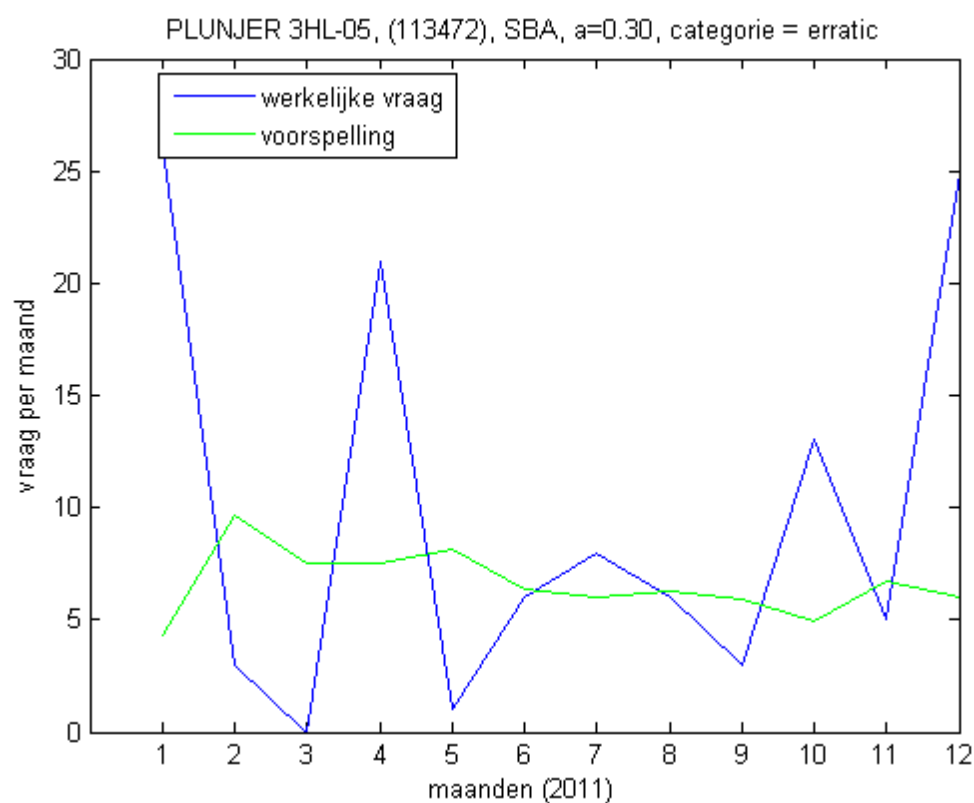
Bijlage 7 – Grafische weergave per categorie

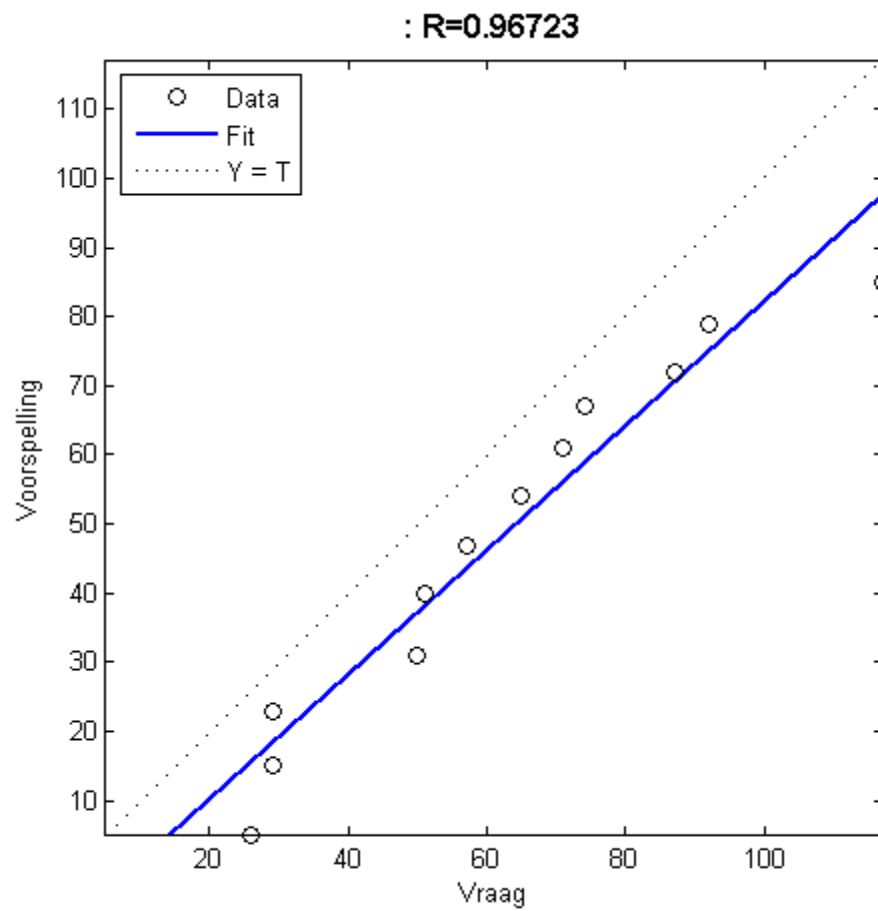
Categorie: Smooth, Voorspelmethode: Croston



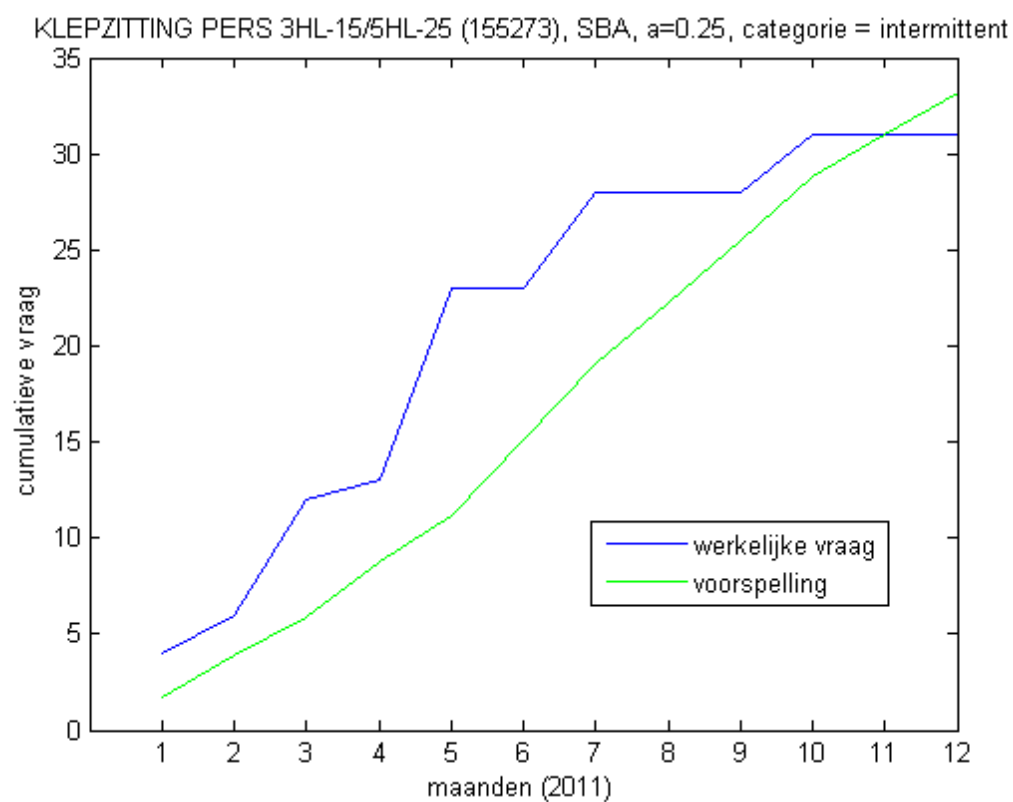
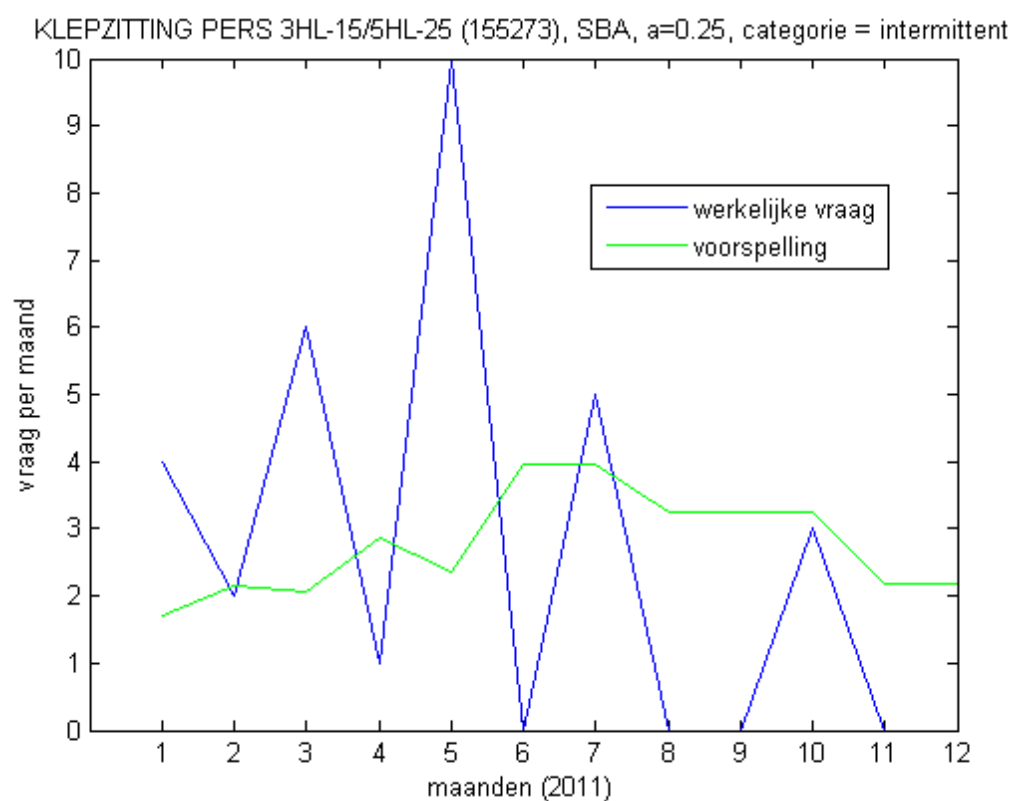


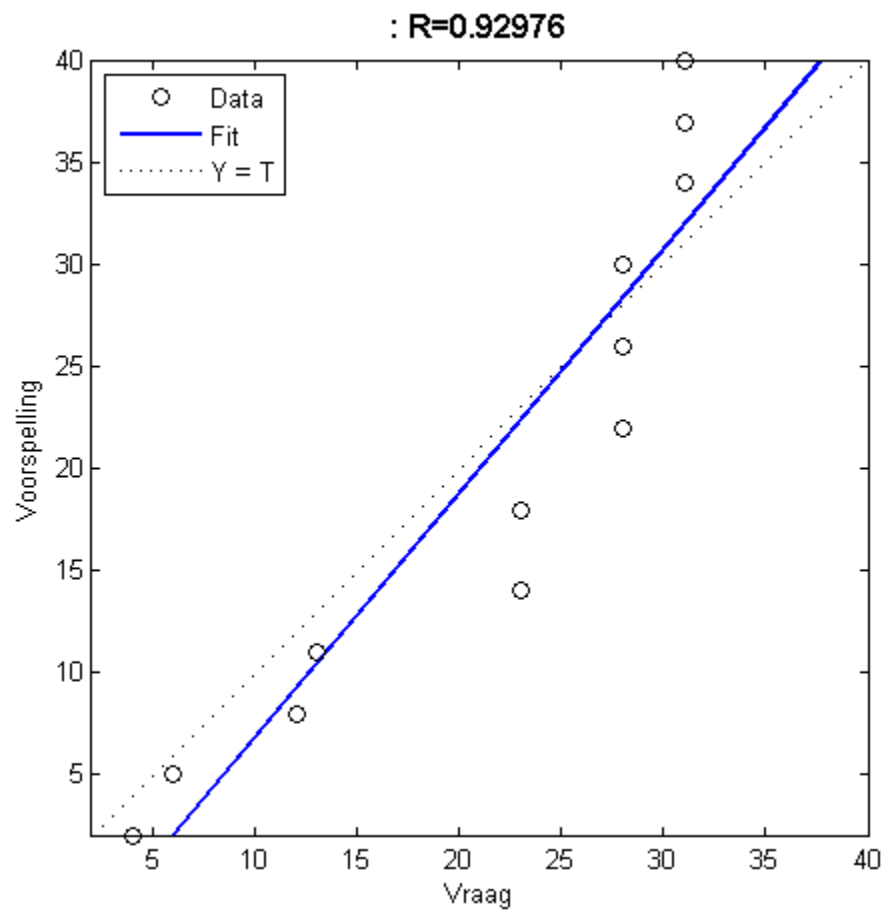
Categorie: Erratic, voorspelmethode: SBA



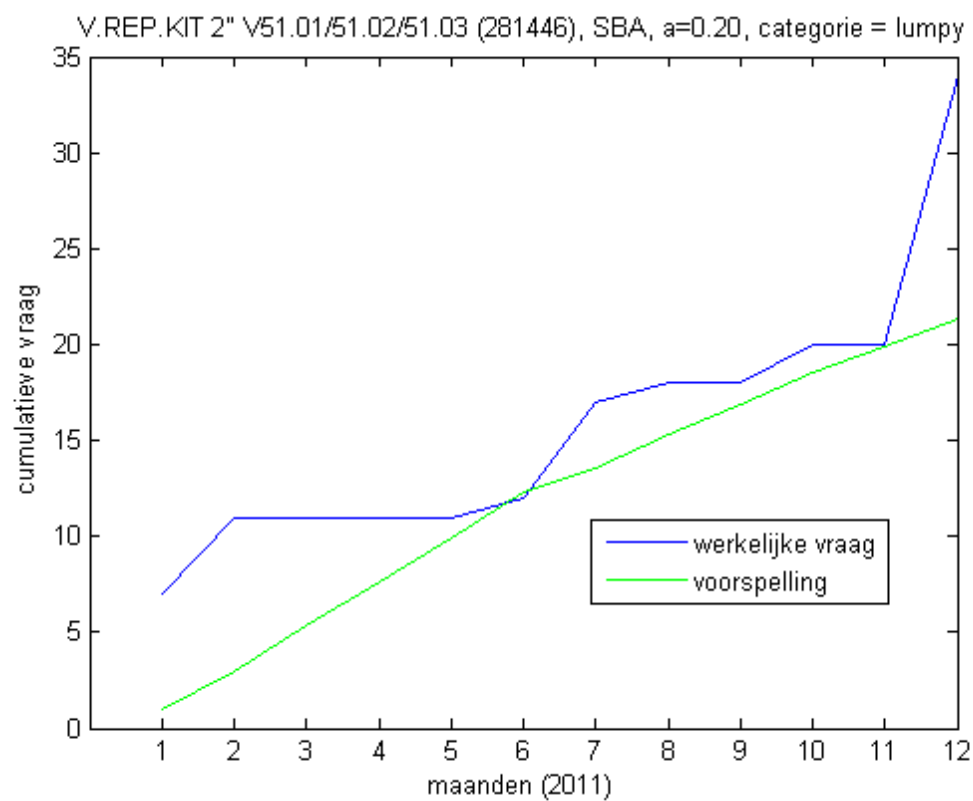
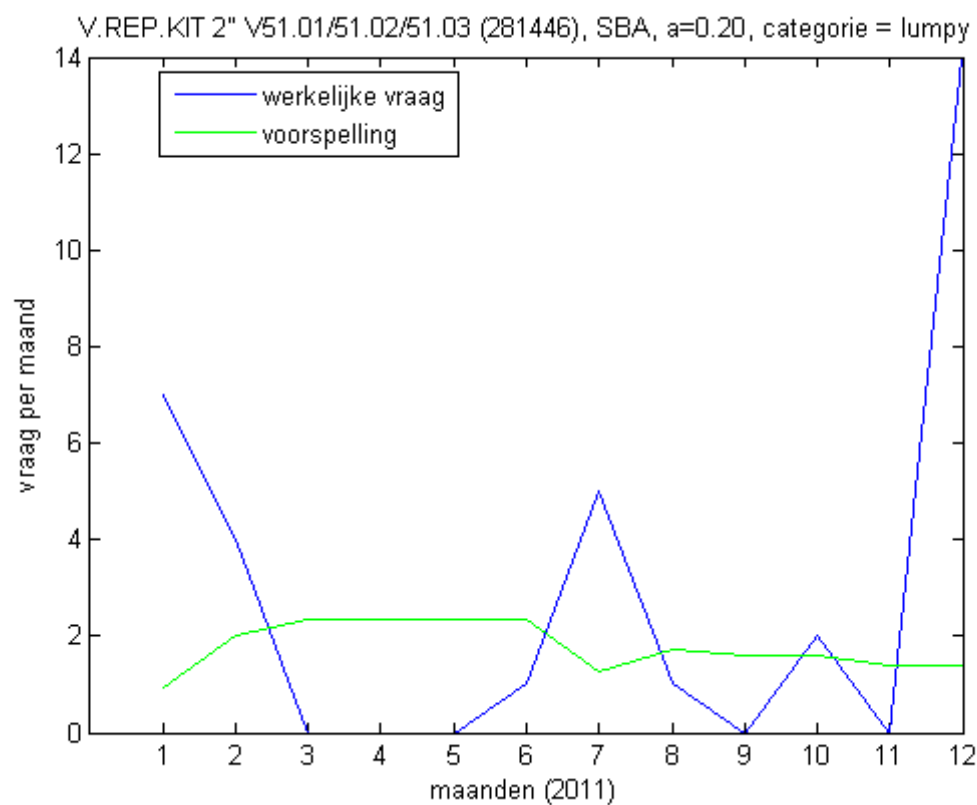


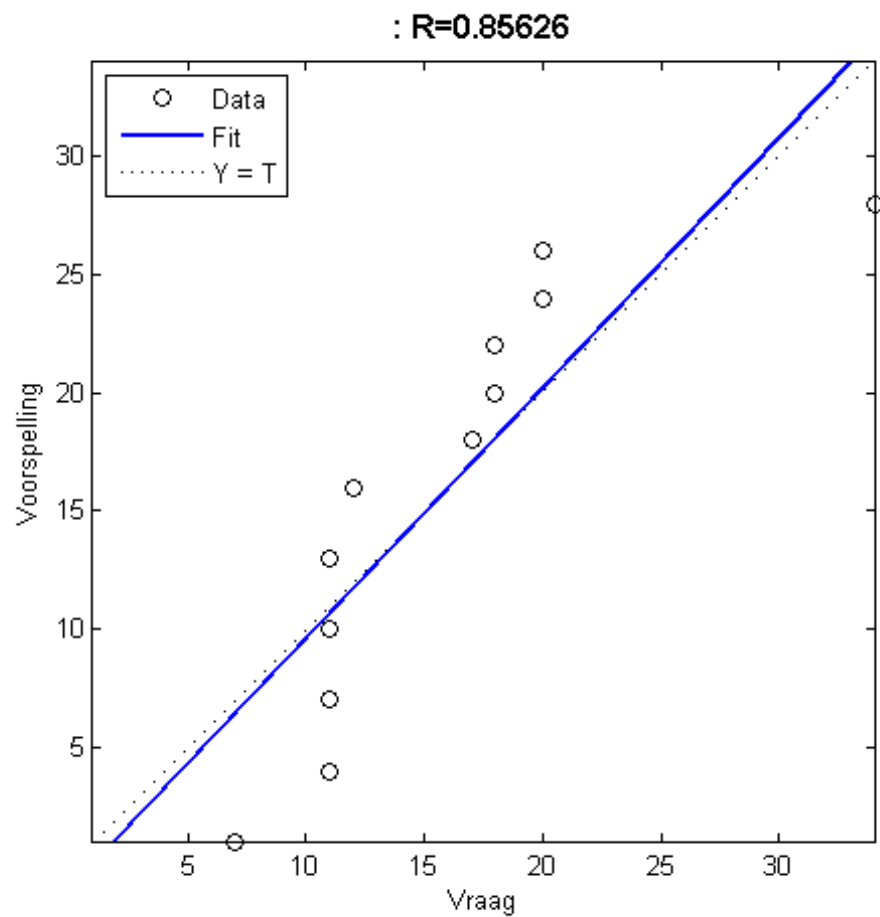
Categorie: Intermittent, voorspelmethode: SBA





Categorie: Lumpy, voorspelmethode: SBA





Classificatie: 1 = smooth, 2 = erratic, 3 = lumpy, 4 = intermittent

Bijlage 8 – Overzicht resultaten van alle artikelen

Artikel-nummer	Beschrijving	Classificatie	Methode	MSE	A-MAPE (%)	CFE	CFE max	CFE min	NOSp (%)	PIS	R	Rico	Start getal
665	FILTERELEMENT	1	Croston	41.667	83.33	-6	7	-22	8.33	105	0.967	1.263	-4
672	KRUISHOOFDPEN 3HL-10/5HL-15	4	SBA	51.333	89.74	20	28	6	75.00	-202	0.962	0.830	-8
672	KRUISHOOFDPEN 3HL-10/5HL-15	4	MA	47.083	93.59	-13	25	-13	58.33	-58	0.954	1.296	-20
1056	GELEIDEBUS	2	SBA	18.917	83.33	5	13	4	83.33	-83	0.985	1.153	-11
1064	LAGERBUS	3	MA	38.167	129.73	-2	11	-5	16.67	23	0.900	1.202	-3
1064	LAGERBUS	3	SBA	29.000	118.92	0	14	-4	8.33	3	0.920	1.394	-10
101531	PLUNJER	4	MA	3.500	153.85	-6	3	-8	16.67	38	0.883	1.938	-5
101531	PLUNJER	4	SBA	3.83	153.85	-8	4	-9	16.67	42	0.908	2.357	-8
106742	PLUNJER	2	SBA	143.083	74.05	53	62	-2	75.00	-475	0.975	0.578	-4
106795	PLUNJER 3HL-10/5HL-15	1	Croston	158.667	53.72	86	86	-10	91.67	-640	0.988	0.621	-1
113472	PLUNJER 3HL-05	2	SBA	102.333	78.63	32	32	6	91.67	-169	0.967	0.899	-8
131774	MANCHET D80/60,9 X 12,7	2	SBA	529.917	76.31	43	46	5	100.00	-303	0.980	1.001	-25
132604	SHAMBAN SEAL PD.13	1	Croston	26916.917	66.45	167	402	-26	83.33	-1670	0.972	0.960	-88
134780	ONTLUCHTER TYPE 11-AV Z	3	MA	7.000	114.29	-8	-2	-14	0.00	106	0.919	1.125	8
134780	ONTLUCHTER TYPE 11-AV Z	3	SBA	5.583	90.48	-3	-1	-9	0.00	60	0.937	0.988	5
135953	PAKKINGSET 55/75	4	SBA	247.000	84.15	36	37	4	75.00	-293	0.974	0.854	-10
135953	PAKKINGSET 55/75	4	MA	199.917	82.32	7	30	-2	75.00	-152	0.980	1.026	-15
137088	MANOMETER 0-400 BAR R.S 250B.	2	SBA	15.000	65.00	16	16	9	91.67	-143	0.967	0.879	-9
137924	BALG V SANITAIRE PULSDEMPER	1	Croston	300.167	17.76	40	82	35	100.00	-724	0.999	0.960	-40
138097	SHAMBAN SEAL PD 13 WM212-T-05	2	SBA	559.917	70.85	117	151	32	91.67	-948	0.986	0.595	-13
138532	MANTELRING TFMA 42,0X34,8X3,8	2	SBA	151.000	74.68	38	46	-27	41.67	-73	0.969	0.587	23

Classificatie: 1 = smooth, 2 = erratic, 3 = lumpy, 4 = intermittent

Artikel-nummer	Beschrijving	Classificatie	Methode	MSE	A-MAPE (%)	CFE	CFE max	CFE min	NOSp (%)	PIS	R	Rico	Start getal
140186	KLEP ALS SERV.	2	SBA	330.500	71.88	42	62	7	91.67	-277	0.964	0.747	2
145939	REPAIRKIT OVERSTR.VENT.NW 25	1	Croston	43.250	93.06	-7	13	-12	16.67	34	0.981	1.215	-6
155270	BEUGEL	2	SBA	1124.417	95.20	27	79	16	75.00	-587	0.980	1.090	-68
155273	KLEPZITTING PERS 3HL-15/5HL-25	4	MA	12.417	112.90	-9	10	-9	50.00	-16	0.929	1.215	-6
155273	KLEPZITTING PERS 3HL-15/5HL-25	4	SBA	11.250	112.90	-9	9	-9	50.00	-11	0.930	1.195	-5
155275	KLEPZITTING ZUIG	4	MA	6.500	112.00	-14	3	-14	33.33	41	0.947	1.579	-7
155275	KLEPZITTING ZUIG	4	SBA	4.417	84.00	1	9	1	66.67	-56	0.952	0.960	-4
155351	LAGERBUS	4	MA	64.250	92.05	-9	33	-9	66.67	-139	0.935	1.226	-25
155351	LAGERBUS	4	SBA	57.417	85.23	5	29	5	75.00	-127	0.965	1.088	-16
155370	KLEP ALS SERV.	2	SBA	87.083	70.19	17	34	1	91.67	-144	0.977	0.776	1
155492	GRONDBUS 70/51/50	1	Croston	761.333	74.07	-6	54	-11	58.33	-124	0.984	1.136	-36
166453	OPT100 ELEMENT 4 DRAADS	1	Croston	69.333	31.37	-2	24	-13	58.33	-50	0.993	1.063	-13
224807	REPAIRKIT	1	Croston	34.417	75.31	-3	16	-8	75.00	-75	0.979	1.259	-20
228091	MEMBRAAN	2	SBA	32.667	75.76	14	21	9	91.67	-166	0.967	0.928	-11
228093	O-RING 60X3 EPDM	2	SBA	993.750	77.40	-73	46	-74	33.33	157	0.974	1.236	-45
235778	O-RING 368 ARP D=196,22X5,33	3	MA	38.167	105.88	-18	6	-18	25.00	67	0.932	1.081	3
235778	O-RING 368 ARP D=196,22X5,33	3	SBA	37.833	90.20	-4	11	-9	50.00	-21	0.934	0.868	2
281446	V.REP.KIT 2" V51.01/51.02/51.03	3	MA	17.417	102.94	9	9	-3	41.67	-32	0.893	0.934	-2
281446	V.REP.KIT 2" V51.01/51.02/51.03	3	SBA	19.167	111.76	6	7	-6	25.00	-1	0.856	1.058	-1
287731	REPAIRKIT TOTAAL NBS701 1 + 1.1/2"	3	MA	11.167	121.43	-12	5	-12	33.33	38	0.946	1.532	-6
287731	REPAIRKIT TOTAAL NBS701 1 + 1.1/2"	3	SBA	10.917	110.71	-5	5	-6	33.33	17	0.953	1.230	-3

Classificatie: 1 = smooth, 2 = erratic, 3 = lumpy, 4 = intermittent

Artikel-nummer	Beschrijving	Classificatie	Methode	MSE	A-MAPE (%)	CFE	CFE max	CFE min	NOSp (%)	PIS	R	Rico	Start getal
298217	ZITTING 1" V	2	SBA	297.667	55.73	46	61	26	100.00	-519	0.988	1.006	-44
306871	REPAIRKIT PRODUCT-PARTS VALVE N701	3	MA	10.333	213.33	-18	-4	-18	0.00	131	0.887	1.360	8
306871	REPAIRKIT PRODUCT-PARTS VALVE N701	3	SBA	6.250	153.33	-9	2	-9	16.67	44	0.885	1.088	3
310325	BALG	2	SBA	16.000	58.46	-2	15	-2	66.67	-57	0.985	1.262	-16
310796	O-RING EPDM FDA-APPROVED 200*8	3	MA	17.167	173.91	-8	3	-8	25.00	31	0.878	0.952	3
310796	O-RING EPDM FDA-APPROVED 200*8	3	SBA	11.417	152.17	-7	3	-12	8.33	58	0.907	1.051	4
A0371160	BINNENBORGRING D60 DIN 472-ST	3	MA	50.417	164.29	-3	19	-3	25.00	-105	0.927	1.379	-21
A0371160	BINNENBORGRING D60 DIN 472-ST	3	SBA	74.417	178.57	-17	21	-17	8.33	92	0.877	1.845	-20
A2156741	PAKKINGSNOER FLOWTITE 10MM L=	1	Croston	465.417	60.56	-5	45	-6	75.00	-162	0.991	0.991	-12
A2163488	MANCHET D55 TYPE N.55-103	1	Croston	6963.500	43.53	46	183	-3	91.67	-1044	0.996	1.007	-95
A2163490	MANCHET D47 TYPE N.47-100	1	Croston	752.583	56.42	-65	32	-100	8.33	583	0.974	1.175	5
A2172292	O-RING 28,25X2,62 SH80 55918PC	3	MA	2391.417	119.74	131	131	-34	16.67	77	0.744	0.378	39
A2172292	O-RING 28,25X2,62 SH80 55918PC	3	SBA	2210.083	112.02	137	137	-28	25.00	-7	0.799	0.357	33
A2172332	O-RING 37,69X3,53 SH80 55918PC	1	Croston	167383.583	48.08	-	207	-3823	8.33	21490	0.998	1.525	-960
A2172334	O-RING 44,04X3,53 SH80 55918PC	1	Croston	72837.167	35.64	194	634	-20	91.67	-2914	0.997	1.004	-259
A2172336	O-RING 50,39X3,53 SH80 55918PC	1	Croston	105665.083	46.26	-41	544	-319	66.67	-1145	0.993	1.017	-159

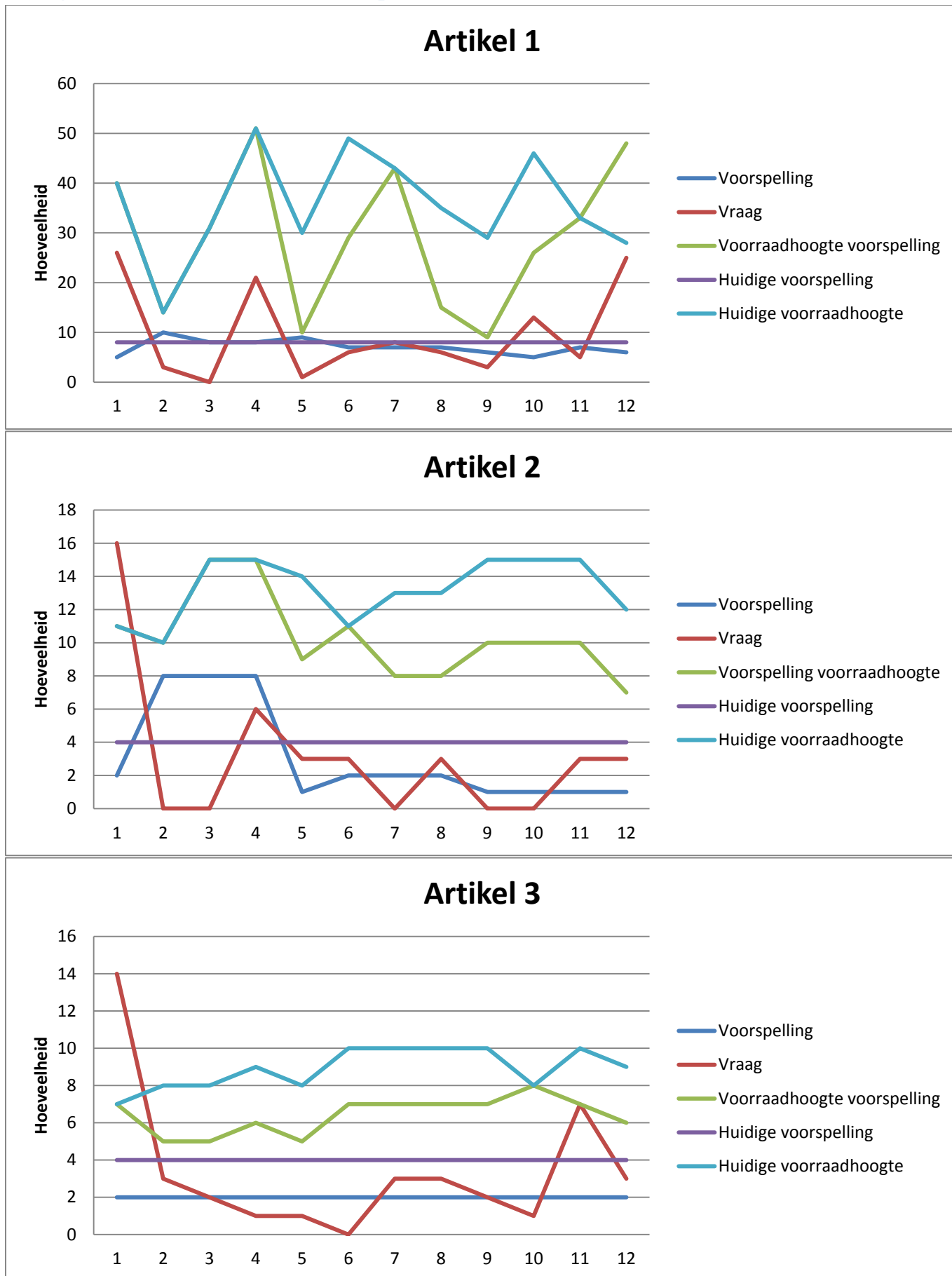
Classificatie: 1 = smooth, 2 = erratic, 3 = lumpy, 4 = intermittent

Artikel-nummer	Beschrijving	Classificatie	Methode	MSE	A-MAPE (%)	CFE	CFE max	CFE min	NOSp (%)	PIS	R	Rico	Start getal
A2172468	O-RING 55,56X3,53 SH80 55918PC	2	SBA	6.083	67.74	7	9	4	83.33	-79	0.987	0.892	-4
A2172901	O-RING 94,92X2,62 VITON 70	3	MA	38.250	167.86	1	19	0	25.00	-75	0.837	0.667	-1
A2172901	O-RING 94,92X2,62 VITON 70	3	SBA	30.750	153.57	1	7	-10	25.00	-6	0.888	0.722	4
A2175115	DICHTINGSRING D50 EPDM 70SH	3	MA	113.083	205.56	-81	-12	-86	0.00	728	0.949	2.426	15
A2175115	DICHTINGSRING D50 EPDM 70SH	3	SBA	68.750	157.41	-13	18	-20	16.67	71	0.895	1.321	-4
A2175117	DICHTINGSRING D65 EPDM 70SH	3	MA	146.083	167.24	-21	-6	-45	0.00	300	0.795	0.845	28
A2175117	DICHTINGSRING D65 EPDM 70SH	3	SBA	120.833	141.38	0	11	-24	16.67	102	0.865	0.673	16
A2180306	O-RING 60,3X58,5X5 EPDM 80	3	MA	219.417	142.25	13	30	-23	16.67	67	0.811	0.416	20
A2180306	O-RING 60,3X58,5X5 EPDM 80	3	SBA	204.083	128.17	7	18	-29	16.67	133	0.801	0.510	23
P0595447	SLIJTBUS 3HL-10	4	MA	0.750	150.00	-5	0	-6	0.00	30	0.891	1.691	0
P0595447	SLIJTBUS 3HL-10	4	SBA	0.667	133.33	-6	0	-6	0.00	31	0.888	1.747	0
P0620794	KLEPZITTING 3HL-10/5HL-15 ARMCO	1	Croston	865.417	53.45	-35	45	-35	33.33	16	0.991	0.991	4
P0632339	DRUKTRANSMITTER	4	SBA	2.417	89.47	-1	1	-6	25.00	18	0.934	1.156	0
P0632339	DRUKTRANSMITTER	4	MA	1.833	63.16	0	1	-5	8.33	18	0.929	1.045	1
P0632722	NIVEAU ZENDER	4	SBA	0.583	100.00	-5	0	-5	0.00	35	0.977	1.783	0
P0632722	NIVEAU ZENDER	4	MA	0.583	100.00	-5	0	-5	0.00	35	0.977	1.783	0
P0678646	SLIJTRING D=42 XMS010 5HL- 25	4	MA	2.583	141.67	-5	4	-5	16.67	17	0.895	1.603	-3
P0678646	SLIJTRING D=42 XMS010 5HL- 25	4	SBA	1.750	84.62	-3	3	-4	16.67	19	0.943	1.455	-2

Classificatie: 1 = smooth, 2 = erratic, 3 = lumpy, 4 = intermittent

Artikel-nummer	Beschrijving	Classificatie	Methode	MSE	A-MAPE (%)	CFE	CFE max	CFE min	NOSp (%)	PIS	R	Rico	Start getal
P0678647	HOMOG.BED XMS010 5HL-25 D=42	4	MA	1.250	69.23	-1	3	-1	33.33	-9	0.962	1.222	-2
P0678647	HOMOG.BED XMS010 5HL-25 D=42	4	SBA	1.750	84.62	-3	3	-4	16.67	19	0.943	1.455	-2
P0694649	STANDMELDKOP UNV.TYPE N0900	4	MA	1.500	85.71	0	1	-3	16.67	7	0.977	0.832	2
P0694649	STANDMELDKOP UNV.TYPE N0900	4	SBA	1.500	85.71	2	2	-3	25.00	4	0.971	0.741	2
P0699369	FILTERELEMENT	4	SBA	13.000	125.00	2	12	-1	41.67	-46	0.936	0.840	-1
P0699369	FILTERELEMENT	4	MA	11.583	121.88	1	9	-2	41.67	-21	0.950	0.863	1
P0706142	DRUKTRANSMITTER 0-100 BAR	2	SBA	6.083	67.74	7	9	4	83.33	-79	0.987	0.892	-4
P0706143	KAST V.DIGITALE UITL.0-400 BAR	3	SBA	8.833	92.31	2	9	-1	58.33	-39	0.949	1.477	-11
P0706143	KAST V.DIGITALE UITL.0-400 BAR	3	MA	8.167	100.00	-2	8	-5	25.00	3	0.934	1.606	-10
P0856348	FILTER H14 SOFILAIR 1560.02.14 (1DS=1ST)	3	MA	7.333	233.33	-14	-4	-15	0.00	133	0.858	1.025	11
P0856348	FILTER H14 SOFILAIR 1560.02.14 (1DS=1ST)	3	SBA	2.333	116.67	-2	1	-5	8.33	24	0.951	0.796	3
P0857144	CARDANAS	4	SBA	6.750	127.78	5	7	-4	25.00	-8	0.953	0.504	3
P0857144	CARDANAS	4	MA	6.500	122.22	4	6	-4	25.00	-6	0.957	0.543	3
P0870651	PAKKINGSET	2	SBA	409.667	60.63	-4	33	-54	41.67	90	0.983	0.830	36
PB704724	GELEIDBAARHEIDS NIVEAU ELEKTRODE	4	MA	2.917	141.67	-7	1	-7	16.67	23	0.944	1.242	0
PB704724	GELEIDBAARHEIDS NIVEAU ELEKTRODE	4	SBA	2.167	116.67	0	3	-2	25.00	-5	0.966	0.747	1

Bijlage 9 – Voorbeelden invloed op voorraadbeheer



Bijlage 10 – Lijst met figuren en tabellen

Figuur 1 – Organisatiestructuur FD Systems.....	1
Figuur 2 – Fasen van voorspelling (uit Boylan & Syntetos, 2009).....	9
Figuur 3 – Classificatie van vraagpatronen	10
Figuur 4 – Superieure voorspelmethodek per kwadrant volgens Syntetos & Boylan	12
Figuur 5 – Superieure voorspelmethodek volgens Kostenko & Hyndman	13
Figuur 6 – Flowchart data-analyse.....	17
Figuur 7 – Aantal artikelen dat in de beschouwde jaren hetzelfde geclassificeerd bleven.....	19
Figuur 8 – Resultaten orders inleggen	33
Tabel 1 – Gekozen prestatie-indicatoren.....	5
Tabel 2 – Prestatie-indicatoren huidige voorspelling	8
Tabel 3 – Gekozen methodek per categorie	14
Tabel 4 – Aantal artikelen per categorie (data 2009-2011).....	17
Tabel 5 – Aantal artikelen dat in de beschouwde jaren hetzelfde geclassificeerd bleven	19
Tabel 6 – Dataset voor verdere analyse (N=239).....	20
Tabel 7 – Aantal te onderzoeken artikelen per categorie	22
Tabel 8 – Voorbeeld smooth: prestatie-indicatoren.....	23
Tabel 9 – Resultaten categorie smooth	24
Tabel 10 – Voorbeeld erratic: prestatie-indicatoren	25
Tabel 11 – Resultaten categorie erratic	25
Tabel 12 – Resultaten Wilcoxon Signed Rank test voor SBA en MA in de categorie intermittent	26
Tabel 13 – Voorbeeld intermittent: prestatie-indicatoren.....	26
Tabel 14 – Resultaten categorie intermittent.....	27
Tabel 15 – Resultaten Wilcoxon Signed Rank test voor SBA en MA in de categorie lumpy.....	27
Tabel 16 – Voorbeeld lumpy: prestatie-indicatoren.....	28
Tabel 17 – Resultaten categorie lumpy	28
Tabel 18 – Gemiddelde waardes van de prestatie-indicatoren voor alle categorieën.....	29
Tabel 19 – Resultaten Wilcoxon Signed Rank Test voor huidige voorspelling met nieuwe voorspelling.....	30
Tabel 20 – Bekende vraaghoeveelheid per maand.....	32
Tabel 21 – Vergelijking tussen de normale voorspelling en de aangepaste voorspelling	33
Tabel 22 - Invloed verschil tussen huidige voorspelling en voorspelmethodek op voorraadbeheer	35

Leeswijzer

ADI: Average Demand Interval – gemiddelde tijd (in maanden) tussen twee vraagmomenten

A-MAPE: Adjusted Mean Absolute Percentage Error - prestatie-indicator die aangeeft hoeveel procent de gemiddelde afwijking is ten opzichte van de gemiddelde vraag.

CFE: Cumulated Forecast Error – prestatie-indicator die de cumulatieve voorspelfout weergeeft.

CR: methodiek van Croston

CV²: Squared Coefficient of Variation – meet de gekwadrateerde verhouding tussen standaardafwijking en gemiddelde.

Erratic: categorie met artikelen die een lage ADI hebben en een hoge CV².

FSS: Forecast Support System

GM: Grey prediction model

Homogeniteit: artikelen die homogeen zijn, zijn artikelen die niet van categorie wisselen over het aantal beschouwde jaren.

Intermittent: categorie met artikelen die een hoge ADI hebben en een lage CV².

Lumpy: categorie met artikelen die een hoge ADI hebben en een hoge CV².

MA: Moving Average

MSE: Mean Squared Error – prestatie-indicator die de gemiddelde gekwadrateerde fout geeft.

NN: Neural Network

NOSp: Number of Shortages – prestatie-indicator die aangeeft hoeveel procent van de maanden er een tekort was.

PIS: Periods in Stock – prestatie-indicator die aangeeft hoeveel en hoelang artikelen er in fictieve voorraad liggen/tekort komen aan het eind van de horizon.

Post-Processing: onderdeel van het FSS dat de voorspelling aanpast aan de hand van andere informatiebronnen.

Pre-processing: onderdeel van het FSS dat de spare parts categoriseert.

Processing: onderdeel van het FSS dat de juiste voorspelmethode selecteert per categorie.

Regressiecoëfficiënt: geeft aan hoe goed er een lineair verband is tussen verschillende datapunten.

SBA: Syntetos & Boylan Approximation

SES: Single Exponential Smoothing

SFDS: Stork Food & Dairy Systems

SKU: Stock Keeping Unit

Smooth: categorie met artikelen die een lage ADI hebben en een lage CV².