

UNIVERSITY OF TWENTE.



MASTER THESIS

Life Cycle Prediction for spare parts at IBM Spare Parts Operations

Ing. Lianne Bensing
Student ID: s1015699

Amsterdam
7/11/2012

Supervisors University of Twente
Dr. A. Al Hanbali
Dr. M.C. van der Heijden

Supervisors IBM
Drs. L.J.H. Neomagus
Ir. J.P. Hazewinkel M.B.A.

Management Summary

Introduction

The research is executed at one of the departments of IBM SPO, being Life Cycle Planning. This department is, amongst other things, responsible for calculating the need to cover the usage of a Field Replaceable Unit (FRU) over the remaining service period (RSP) in case a supplier stops producing and IBM has a final opportunity to acquire parts, known as a Last Time Buy (LTB). Due to the time horizons, that can be up to 16 years, the amount of parts to acquire is difficult to forecast.

Motivation

Based on a pilot project executed for the division Lenovo, focused on Commodity Based Lifecycle Forecasting (CBLF), the idea existed that there should be a commonality between the FRUs with a similar usage pattern, which can be used for clustering. Having the opportunity to assign an FRU to a cluster would make long term forecasting easier, since the usage pattern is known. Therefore the aim of this research is the following:

“Investigate which characteristics of an FRU are related to a specific usage pattern and how this information can be used to cluster FRUs into groups with a similar usage pattern, in order to improve the forecast accuracy.”

Research methodology

The start of the research was the determination of the current forecast accuracy and the different types of usage patterns the FRUs follow. After the FRUs were assigned to a specific usage pattern, an investigation of the relation between the usage pattern and a set of characteristics was executed, to determine whether characteristics can be related to a specific usage pattern and can be used to cluster the FRUs. Therefore we used a combination of statistical testing, data analysis and factorial analysis. As a second step, we assessed the performance of the pilot that triggered the research, and we discussed and tested options to improve this method, based on a simulation of historical LTBs.

Results

Based on the analysis of the usage patterns, 12 different partial usage patterns were identified over a period of 5 year historical usage, because that is the amount of years for which historical data is stored. It appeared to be impossible to combine these partial usage patterns into a limited set of usage patterns, due to the large amount of possible combinations. With respect to the forecast accuracy, the performance is determined based on the bias, Mean Absolute Deviation (MAD) and the Mean Absolute Percentage Error (MAPE). The results indicated that the standard decline approach, in which a fixed decline factor is used for every year the need has to be forecasted, has an aggregated MAPE value of 235% against 314% for CBLF. As a result, standard decline leads to more accurate forecasts on an aggregated level. CBLF has a more accurate result when it is actually the best approach, with an average MAPE of 203% for CBLF compared to an average MAPE of 241% for the standard decline. However, the results of both approaches indicate room for improvement, because the average aggregated difference between the forecasted and actual need is more than 50%, which indicates an inaccurate forecast.

With respect to identification and clustering we investigated the possibilities for 8 characteristics, being Age FRU, Brand, Commodity, Division, Forecasted Reliability, LTB Month forecast, RSP and

TM144 status (indicates if the LTB is executed before or after production stops). On an aggregated level, statistical tests indicate that all characteristics could have a relationship with a partial usage pattern. On commodity level, with a commodity being a group of FRUs with similar purpose (like batteries), some of the characteristics can be excluded based on the statistical tests. We focused on clustering approaches for the commodity HDD, based on a combination of ranges, sub ranges and historical usage data. The amount of FRUs selected using the clustering approaches ranges in most cases between 1 and 5 FRUs, resembling less than 5% of the entire selection. The percentage of FRUs successfully identified can be high, but is not considered to be representative, since 1 FRU correctly identified from a selection of 1 can also be coincidence.

In the second part of the research, we focused on the performance of CBLF. The analysis pointed out three main causes for an inaccurate forecast with CBLF, being a difference between the observed usage pattern and the pattern used to forecast future usage with, a difference between the expected and actual time on the market and a difference between the point in time the observed usage starts and the point in time in which the usage starts based on the curve used to forecast future usage. Different improvement options are tested to minimize the effects of the main causes and all options have realized accurate forecasts for at least 1 FRU in the selection, but a straightforward method to determine the most appropriate forecast method for all FRUs could not be developed.

Of the different forecasting options considered, the aggregated MAPE values for a selection of 13 fast moving FRUs are 63% for Weibull, 71% for Gamma and 133% for standard decline. For slow movers, the aggregated MAPE values for a selection of 6 FRUs are 78% for scaled CBLF, 114% for Croston's method and 139% for standard decline. These methods are combined in an Excel tool, that can be used to visualize the possible usage patterns an FRU might follow and indicate what ranges of usage might be realized. This could help the SPO team in making a decision about the amount to acquire, but can also clarify the advantages that range forecasting can offer.

Main conclusions and recommendations

Based on this research, the main conclusions are:

- Both characteristic values and the most appropriate forecast method are FRU specific.
- Characteristic values can give an indication of the usage pattern.
- Clustering approaches cannot be determined based on partial usage data. Full life cycle data might provide additional insights that makes clustering possible.
- In the process of determining the LTB need, it is vital to determine the most appropriate forecasting method based on data of the specific FRU.
- Promising forecasting methods for fast movers are Gamma, standard decline and Weibull. For slow movers, promising forecasting methods are scaled CBLF, standard decline and Croston's method.

To improve the forecast accuracy, we recommend:

- Start using forecasting methods that are specifically designed for slow movers.
- Focus on acquiring the necessary data for range forecasting, to create better insight in the possible future needs, and expand the time period over which historical usage data is stored.
- Apply the Excel tool with the different forecast methods in the time required to get the data regarding range forecasting, to get acquainted with range forecasting and the possible benefits it could offer.

Preface

"You can never plan the future by the past"

Edmund Burke (1729 – 1797)

Even though Edmund Burke already stated somewhere in the 18th century that we cannot plan the future by the past, it is still something we try today. Planning the future by the past is namely the main theme of my thesis, and it will stay the main theme of many other projects in the future as well. In my case, the project is finished, but the statement that has proven to be true during the process is:

"Time is the wisest counselor of all"

Pericles (495 – 429 BC)

Although time may be the wisest counselor, it is not the only one that helped me throughout the process of researching and writing my thesis. For that, many other people deserve the credits, and I am grateful they all spent time to help me.

First I would like to thank the team of IBM SPO, for giving me the opportunity to do the research, for making me feel welcome and for the willingness to answer my questions and providing help and advice when necessary. I really enjoyed being at the office and I learned a lot about all kinds of subjects, not only related to forecasting and SPO processes, but also on a personal level. Special thanks to Laurens, for his guidance in this process and the good conversations. Furthermore, I would like to thank Daniëlle, Corine, Hans, Ron, Jaap, Cor, Menno, Dennis and everyone I might forget for their help and interest in my research, I really enjoyed working with you. Finally, I would like to thank Johan for the fun conversations during the day, which could really help lighten the process when my research did not go according to plan.

I want to thank Ahmad and Matthieu for supervising me from the side of the university. The progress meetings were very helpful, guiding me in the right direction at moments I tended to lose sight of where I was heading. Your knowledge of the subject really helped improving my thesis. Sina, thanks for the additional feedback, tips and interest. And thanks to Chen, I really enjoyed working with you on the ProSeLo project.

And last but not least, I would like to thank my friends and family. Thanks to my fellow students and friends André, Erwin, Robert and Martijn for supporting me during my research and for providing interesting solutions to solve the problems I encountered along the way. Thanks to my mom, for your interest in my assignment, supporting me and taking such good care of me. And finally, a special thanks to my love, Welmer, for your love, support and patience during this project, but also in the years before.

Lianne Bensing

Harderwijk, 7 november 2012

Index

Management Summary.....	II
Preface.....	IV
List of abbreviations	VIII
List of definitions	IX
List of figures	XI
List of tables	XII
1 Introduction.....	1
1.1 IBM	1
1.2 IBM Service Parts Operations.....	1
1.3 Products and services.....	2
1.4 Conclusion	2
2 The research.....	3
2.1 Motivation	3
2.2 Research aim	4
2.3 Scope	4
2.4 Research questions.....	4
2.5 What will it bring SPO?	5
2.6 Thesis outline.....	5
2.7 Conclusion	5
3 What is the current situation at the planning department of SPO?	6
3.1 Forecasting processes at SPO	6
3.2 Commodity Based Lifecycle Forecasting	9
3.3 Data issues.....	12
3.4 Case selection.....	13
3.5 Current performance.....	17
3.6 Improvement possibilities	22
3.7 Conclusion	23
4 Model prerequisites and literature background	24
4.1 What are the prerequisites for applicability to the situation of SPO?	24
4.2 The basic PLC forecasting concept	24
4.3 What literature is available regarding PLC clustering?	25
4.4 Which PLC forecasting methods are described in literature?.....	26
4.5 Which methods will be applied at SPO?.....	26

4.6	Conclusion	27
5	Which possibilities exist regarding product clustering using FRU characteristics?.....	28
5.1	Which characteristics can be promising?	28
5.2	Which characteristics influence the PLC of the FRU?.....	30
5.3	Which characteristics can identify the partial usage patterns?	33
5.4	Can a combination of characteristics identify partial usage patterns?	36
5.5	Can characteristics be used to identify fast and slow moving FRUs?	40
5.6	Can the characteristics identify the usage pattern and PLC stage?	41
5.7	Conclusion	43
6	What is the impact of the results on the current way of working?	44
6.1	What is the performance difference between CBLF and standard decline?	44
6.2	What are possible causes for inaccurate CBLF performances?	47
6.3	Which methods exist to improve the performance of CBLF?	51
6.4	Which of the possible methods is the most promising?	57
6.5	Selection method	58
6.6	Conclusion	59
7	How should IBM implement the results?.....	60
8	What are the effects when the results are implemented for all FRUs?	61
8.1	What is the applicability of the results to other FRUs?	61
8.2	What is the applicability of the results to other processes?	61
8.3	What are the advantages and disadvantages?	62
8.4	Conclusion	62
9	Conclusions, recommendations and further research	63
9.1	Conclusions.....	63
9.2	Recommendations.....	63
9.3	Further research	64
	References.....	XIII
	Appendices	XV
A.	Divisions of IBM SPO	XVI
B.	The Product Life Cycle at IBM.....	XVIII
C.	Example effect CBLF curve using average and summed indexed usage	XX
D.	Usage pattern explanation	XXIV
E.	Explanation commodities.....	XXVI
F.	Results statistical tests	XXVII

G.	Outlier analysis	XXXI
H.	Box plots test selection	XXXIII
I.	Description commodities CBLF	XXXV
J.	Standard distributions	XXXVIII

List of abbreviations

CB	Central Buffer
CBLF	Commodity Based Lifecycle Forecasting
EMEA	Europe, Middle East and Africa
EOS	End of Service
FRU	Field Replaceable Unit
GA	General Announcement
IB	Installed Base
IBM	International Business Machines
LTB	Last Time Buy
MAD	Mean Absolute Deviation
MAPE	Mean Absolute Percentage Error
MSE	Mean Squared Error
MVS	Multi Vendor Services
PAL	Parts Availability Level
PDT	Parts Delivery Time
PLC	Product Life Cycle
PLCM	Product Life Cycle Management
RSP	Remaining Service Period
RSS	Retail Storage Systems
SPO	Spare Parts Operations
TCO	Total Cost of Ownership

List of definitions

Characteristics

FRU-related information that is known at the moment of the LTB, and can be used to assess the type of usage pattern. An example of a characteristic is size, because the size of an FRU is known from the moment of introduction onwards.

Commodity

A set of FRUs that have the same characteristics and might also be mutually exclusive. An example of a commodity is batteries, which are all used to power a machine.

Commodity Based Lifecycle Forecasting

A forecasting method based on the assumption that the FRUs in a specific commodity all follow a standard usage pattern, that can be used to forecast the usage over the remaining service period.

End of Service

The moment at which IBM announces not to service a specific machine or FRU anymore.

Field Replaceable Unit

A spare part of IBM. An FRU can consist of multiple parts, that together form a replaceable unit that can be used to repair the machine of a customer.

Installed base

The amount of machines installed at customers at a specific point in the Product Life Cycle.

Last Time Buy

The ability to buy one final amount of parts at the moment a supplier announces to stop the production of that specific part.

Partial usage pattern

The usage pattern of a selection of the Product Life Cycle of an FRU.

Usage pattern

The usage pattern of an FRU over the entire Product Life Cycle.

Product Life Cycle

The time between the moment the FRU is introduced and the moment the FRU is not available for the customer any more.

Remaining Service Period

The amount of years between the moment the LTB for a specific FRU is executed and the moment that FRU will become EOS. The longer the remaining service period, the more difficult it is to forecast the amount of FRUs required.

Spare part

A part of a machine that can be replaced at the moment the part that was originally in the machine breaks down. An example of a spare part is a hard disk, for which a new one can be installed after the previous hard disk has broken down.

TM144 Status

This status is used in the LTB process and can be either PRE or POST. A PRE status indicates that the LTB takes place before production stops, POST means after production stops. When the status is PRE, Manufacturing is responsible for acquiring the correct amount of FRUs, and SPO only needs to provide information regarding the amount of FRUs they expect to use after production stop. When the status is POST, SPO is responsible.

Usage

Historical orders that were places by customers for a specific FRU.

List of figures

Figure 1-1: Net income of IBM divided per product type (based on (IBM, 2011))	1
Figure 1-2: Organization structure of SPO EMEA.....	1
Figure 2-1: Example of a Product Life Cycle (adapted from (Write A Writing))	3
Figure 3-1: Date of birth determination in PLC forecasting	10
Figure 3-2: Example of the result of CBLF for a hard disk	11
Figure 3-3 α - l Identified usage patterns	17
Figure 3-4: Box plot MAPE values per partial usage pattern	20
Figure 3-5: Box plot relative bias values.....	20
Figure 3-6: Box plot relative MAD values	20
Figure 5-1: Example preferred results Box plots	34
Figure 5-2: Box plot LTB Month forecast.....	35
Figure 5-3: Box plot Age FRU	35
Figure 5-4: Indexed usage pattern based on correlated HDDs	43
Figure 6-1: Example effect reallocating usage	50
Figure 6-2: Difference summed and average indexed usage curves	56
Figure A-1: IBM Mainframe	XVI
Figure A-2: Examples of systems from the Power division	XVI
Figure A-3: Disk storage system on a standalone basis	XVI
Figure A-4: System X processors	XVII
Figure B-1: Stages in a Product Life Cycle at IBM	XVIII
Figure C-1: PLC results for example CBLF	XXI
Figure C-2: Actual and smoothed PLC curve with partial information known	XXI
Figure C-3: Actual and smoothed PLC curve with full information known	XXI
Figure H-1: Box plot LTB Month forecast test selection.....	XXXIII
Figure H-2: Box plot Age FRU test selection.....	XXXIV
Figure I-1: Active backplane	XXXV
Figure I-2: I/O riser card	XXXVI
Figure I-3: Light path	XXXVI
Figure I-4: SCSI adapter	XXXVII
Figure J-1: Examples Weibull distribution.....	XXXVIII
Figure J-2: Examples Gamma distribution.....	XXXVIII
Figure J-3: Examples Beta distribution	XXXVIII

List of tables

Table 3-1: General range overview of the initial selection	15
Table 3-2: Percentage occurrence usage patterns.....	17
Table 3-3: MAPE performance classification according to Chien et al. (2010)	19
Table 3-4: Information partial usage patterns	20
Table 3-5: Current performance standard approach, per commodity	21
Table 5-1: Characteristics and their way of measurement	30
Table 5-2: Characteristic values increasing pattern HDD	36
Table 5-3: Selected sub ranges increasing usage pattern	38
Table 5-4: Results identification of increasing usage pattern with characteristic sub range combinations	38
Table 5-5: Historical usage parts with increasing pattern.....	39
Table 6-1: Aggregated forecast accuracy CBLF and Standard decline	45
Table 6-2: Aggregated forecast accuracy most accurate forecasting method.....	45
Table 6-3: Forecast accuracy slow moving FRUs.....	46
Table 6-4: Forecast accuracy fast moving FRUs	47
Table 6-5: Most accurate forecasting method commodities	48
Table 6-6: Forecast accuracy CBLF and scaled CBLF.....	52
Table 6-7: Most accurate forecasts per parameter estimation method.....	54
Table 6-8: Forecast accuracy standard distributions per parameter estimation method	54
Table 6-9: Forecast accuracy CBLF and shifted CBLF.....	55
Table 6-10: Forecast accuracy different curve creation methods	56
Table 6-11: Number of good forecasts per forecasting method for a selection of 19 FRUs.....	57
Table 6-12: Performance forecasting approach selection methods	59
Table C-1: Partial usage data for example CBLF	XX
Table C-2: Indexed partial usage data for example CBLF	XX
Table C-3: Full usage data for example CBLF.....	XX
Table C-4: Indexed full usage data for example CBLF	XX
Table C-5: Percentages differences partial usage for example CBLF	XXII
Table C-6: Percentage differences full usage for example CBLF	XXII
Table C-7: Expected requirement results for example CBLF.....	XXII
Table C-8: Percentage differences partial usage.....	XXII
Table C-9: Percentage differences full usage	XXIII
Table C-10: Expected requirements average and summed indexed usage	XXIII
Table F-1: Results Chi Squared test of Independence.....	XXVII
Table F-2: Results Kruskal-Wallis	XXVII
Table F-3: Number of cases with a p-value less than 0,05 for characteristic values for Division	XXVIII
Table F-4: Number of cases with a p-value less than 0,05 for characteristic values for Brand	XXIX
Table F-5: Number of cases with a p-value less than 0,05 for characteristic values of TM144 Status	XXIX
Table F-6: Resulting p-values hypothesis testing of characteristics on commodity level using Kruskal-Wallis	XXX

1 Introduction

This master thesis is related to a master assignment, which is executed within IBM. In this section we will give an introduction to IBM and the relevant departments, products and services.

1.1 IBM

IBM, or International Business Machines Corporation, started in 1911 with the fusion of three companies. From the beginning, IBM has focused on products and services related to storing, processing and analyzing information. During the last 100 years of business, IBM introduced revolutionary products, like the electronic calculator and the personal computer.

Income per product type

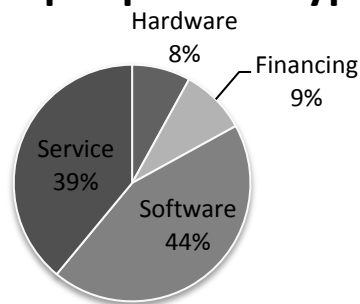


Figure 1-1: Net income of IBM divided per product type (based on (IBM, 2011))

Nowadays, IBM is a well known international player in the business-to-business IT market, providing hardware, financing, software and services to their customers. In the year 2010, IBM operated in 170 countries, employing 426 thousand employees worldwide. The company generated a revenue of \$99.9 billion and a net income of \$14.8 billion in that same year (IBM, 2011). The percentage of the net income generated by the different product types can be found in Figure 1-1.

IBM Netherlands is headquartered in Amsterdam. At this location, operations are carried out for the region Europe, Middle East and Africa (EMEA). One of the departments located in Amsterdam is the Spare Parts Operations (SPO) department, at which the research will be conducted.

1.2 IBM Service Parts Operations

The SPO department is responsible for different aspects with respect to planning, delivery, repair and supporting activities related to spare parts. The assignment is executed at the SPO planning department in Amsterdam. The planning department is responsible for planning, distributing and controlling the spare parts inventory in the Central Buffer (CB) for EMEA, located in Venlo, and in different countries in EMEA.

The department is divided into three sub departments, being Country Planning, Life Cycle Planning and the Planning Control Tower, as can be seen in Figure 1-2. Country Planning is responsible for the spare parts located in the different countries. Life Cycle Planning is concerned with the planning of high-end

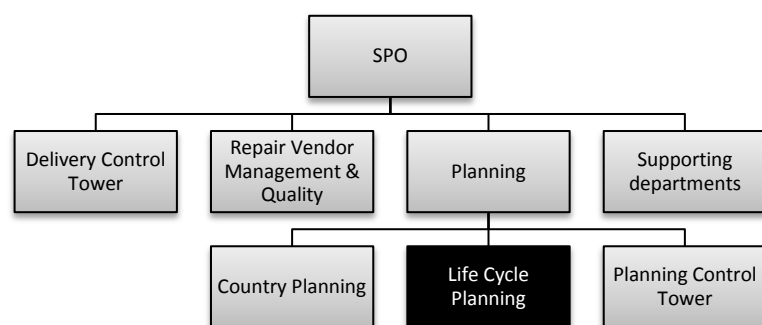


Figure 1-2: Organization structure of SPO EMEA

products, Inventory Management and Product Life Cycle Management (PLCM). The Planning Control Tower is concerned with the planning of low-end products and monitors the planning process, among other activities, to make sure the service targets are reached.

1.3 Products and services

IBM offers its customers different types of machines with different types of service contracts, ranging from 2 hours till next day service. To be able to meet the service requirements, SPO has 277 stock locations in 63 different countries. For every machine IBM offers to its customers, a number of Field Replaceable Units (FRU) are defined by the engineering department. FRUs can for example consist out of a number of spare parts and an instruction form for the Customer Engineer. The FRUs can be stocked at the CB, at a local storage location or not stocked at all.

The products IBM offers are divided over 7 divisions. These divisions are Lenovo, Mainframe, Multi Vendor Systems (MVS), Power, Storage, System X and Retail Storage Systems (RSS). A short description of the divisions can be found in appendix A.

The FRUs can be divided into commodities. Every commodity has its own specific characteristics. However, it is possible that a specific commodity can be present in multiple divisions, like batteries. In that case, characteristics can also differ for FRUs from a specific commodity between different divisions.

1.4 Conclusion

IBM is a worldwide IT company, focused on the business-to-business market. IBM offers hardware, software, financing and services to its customers. Within IBM, the SPO department is responsible for planning, delivery, repair and supporting activities, all related to spare parts, called FRUs. The FRUs belong to a specific commodity and they can be divided over 7 different divisions, being Lenovo, Mainframe, MVS, Power, Storage, System X and RSS.

2 The research

This section focuses on the research, by explaining the motivation, research aim and scope. The research questions will be discussed, followed by the benefits for the company and an outline of the remainder of the thesis.

2.1 Motivation

The ProSeLo Project (Proactive Service Logistics for Advanced Capital Goods), initiated by Dinalog, is a project focused on researching innovative solutions to improve system uptime and the competitive advantage, while reducing the total cost of ownership (TCO) of the product. The ProSeLo project consists of three work packages. IBM participates in work package two, which is focused on Last Time Buy (LTB) and re-use. Other participants in this work package are the University of Twente, Océ Technologies and Vanderlande Industries (Dinalog, 2010).

When a supplier announces an LTB, that supplier stops producing a specific FRU, but IBM might still use the FRU in machines that are produced, sold or serviced. At the moment of the announcement, IBM has the opportunity to procure one last quantity of the FRU from the supplier, that should be sufficient to cover the need until the moment the FRU reaches the end of service (EOS) date. If IBM purchases too much, the FRUs will be scrapped after EOS. If the amount is not sufficient to cover the need in the remaining service period (RSP), penalty cost due to missing contractual service level agreements might occur. With an accurate LTB model, more accurate decisions about the amount of FRUs to procure can be taken, which should reduce the TCO.

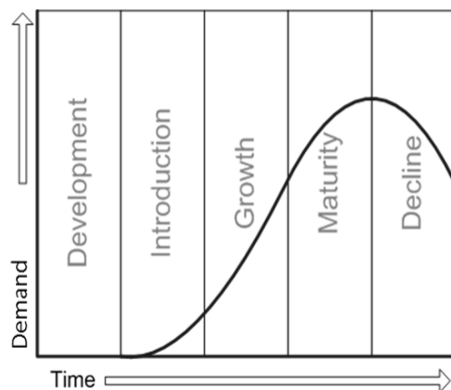


Figure 2-1: Example of a Product Life Cycle (adapted from (Write A Writing))

The accuracy of the forecast in the normal forecasting process and in the LTB model could be improved using Product Life Cycle (PLC) forecasting. A PLC shows the usage pattern of a specific FRU during the period of time the FRU is used by IBM, and is different for every FRU. An example of a PLC can be found in Figure 2-1. PLC forecasting focuses on forecasting the shape of the PLC of an FRU. When this shape is known, more accurate forecasts can be made about the need for the FRU at a specific point in time, which is expected to improve the quality of the LTB decisions.

Within SPO a number of initiatives are initiated to improve the performance of the LTB calculations, with as main goal lowering the costs. One of the initiatives is Commodity Based Lifecycle Forecasting (CBLF), which focuses on implementing the usage pattern of an FRU in the forecasting and planning process. A pilot is currently executed concerning a number of FRUs in the System X division.

In CBLF, only the date of birth of an FRU, the commodity and the actual usage per period are taken into account to generate a PLC curve. However, the SPO team expects that other aspects can be added to the current procedure to fine-tune PLC forecasting even more.

In this research, the assumption is made that PLC forecasting can be based on FRU-specific characteristics. Examples of these characteristics are weight, vitality of the part to the functioning of the machine, the installed base (IB) of the machine and the replenishment rate. The characteristics

that actually influence the PLC might differ per FRU. This research is initiated to investigate the characteristics that influence the PLC of an FRU.

Based on the given description, the following problem statement can be determined:

In order to improve the quality of LTB decisions, IBM would like to incorporate PLC forecasting in their processes. It is currently unclear which characteristics, besides usage, IBM can apply in order to create more accurate PLC forecasts and cluster FRUs to reduce forecasting complexity.

2.2 Research aim

The aim of this research is to distinguish if there are, besides usage, additional characteristics which influence the PLC of an FRU and to what extent they influence the PLC. The results of the research are to be embedded in the forecasting process, to improve PLC forecasting and thus improve the forecast accuracy in the LTB process, and if possible in the normal forecasting process, while reducing the obsolescence costs within SPO.

2.3 Scope

The scope of the research is defined to be able to execute the research within a time frame of six months, which is the general guideline for the length of a master thesis. The aspects that will be incorporated into the research are:

- The research will be carried out for the EMEA region, with a focus on the CB.
- A limited set of FRUs will be selected to be investigated thoroughly. The size of the set of FRUs is yet to be determined, based on a set of selection criteria.
- The entire PLC will be taken into account in the analysis of the selected FRUs. A detailed description of the PLC concept can be found in appendix B.
- Only demand forecasting is incorporated into the scope. Returns and warranty forecasting are not taken into account.
- Actual implementation of the results is not included in the scope, but an implementation plan will be created.
- If improvement possibilities related to aspects outside the research scope are detected, these will be mentioned, but not thoroughly investigated.

2.4 Research questions

Research questions are used to structure both the phases in the research and the chapters to be discusses in the master thesis. The research aim will be achieved by answering the main research question, which is:

How can IBM apply PLC forecasting, using part-related characteristics, such that forecasting accuracy improves?

To answer the main research question, a number of sub questions are specified. These sub questions, and the points of attention related to them, are:

1. What is the current situation at the planning department of SPO?
 - a) Analyze current forecasting process
 - b) How is PLC forecasting currently applied
 - c) Selection and analysis of a set of FRUs for detailed investigation
 - d) Quantify the current performance of the FRUs

- e) Improvement possibilities
- 2. Model prerequisites and literature background
 - a) Prerequisites for applicability to the situation of IBM
 - b) The PLC concept
 - c) Available literature regarding characteristics
 - d) Available literature regarding PLC forecasting
 - e) Applicability
- 3. Which possibilities exist regarding product clustering using FRU characteristics?
 - a) Determine promising characteristics from literature, expert opinions and available data
 - b) Determine the influence of the characteristics
 - c) Investigate ability to identify and cluster on usage patterns
- 4. What is the impact of the results on the current way of working?
 - a) Assess performance difference between forecasting methods
 - b) Determine possible causes for inaccuracies
 - c) Test methods to improve the performance
 - d) Test how to select the most promising method
- 5. How should IBM implement the possible improvements?

In this sub question an implementation plan will be proposed, which can be used by IBM in the implementation of the research results.
- 6. What are the effects when the results are implemented in the forecasting process for all FRUs?
 - a) Assess applicability of the results to other FRUs
 - b) Assess potential when applying to other FRUs
 - c) Advantages and disadvantages

2.5 What will it bring SPO?

The research outcome will give SPO additional insight in possible similarities and dissimilarities between the PLCs of the FRUs and opportunities to cluster FRUs to reduce the amount of potential usage patterns. Furthermore, the research will give input regarding forecasting techniques that might, or might not, be suited to forecast the future need and the improvements that could be realized by changing the forecasting process.

2.6 Thesis outline

Chapter 1 and 2 give an introduction to the company IBM and the research. Chapter 3 describes the current way of working at SPO and the selection of FRUs to investigate. The results of the literature study can be found in chapter 4 and the investigation of clustering based on characteristics in chapter 5. The effects of the improvement options are mentioned in chapter 6. An implementation plan is created in chapter 7 and the possibilities when the results are extrapolated can be found in chapter 8. Conclusions, recommendations and topics for further research are listed in chapter 9.

2.7 Conclusion

The research is related to the ProSeLo Project work package two, focused on Last Time Buy and Re-use. The focus of this research is on improving the forecast using PLC forecasting and the aim is to distinguish if there are, besides usage, additional characteristics which influence the PLC of FRUs, which characteristics it concerns and to what extent they influence the PLC. A main research question and 6 sub questions are formulated to guide the research within the specified scope.

3 What is the current situation at the planning department of SPO?

The starting point of the research is the description of the forecasting processes applied at SPO, which form the framework for the research. Once the processes are known, cases can be selected. Based on these cases we can describe the current situation at SPO with respect to forecasting accuracy, and improvement options can be determined.

3.1 Forecasting processes at SPO

Within SPO, different forecasting processes are executed. Forecasting focuses on day-to-day, new products, LTB and PLC forecasting. The different forecasting processes, with the important characteristics, assumptions and requirements corresponding to the processes, are described in this section. An exception is CBLF, which is described in detail in section 3.2.

3.1.1 Day-to-day forecasting at CB level

Day-to-day forecasting at CB level is focused on forecasting the amount of FRUs that are expected to be used on a short term to serve customers on an aggregated level, i.e. not country specific. The forecast is created by Xelus, the forecasting program of SPO, based on four-week time buckets. To get a forecast, the planner needs to select one or multiple forecasting techniques preprogrammed or manually created in Xelus. For a large percentage of the FRUs, the preprogrammed forecasting techniques moving average, single smoothing or double exponential smoothing are selected. There is no distinction made between the methods selected for slow or fast movers, although in literature this distinction is made, with slow movers being defined as parts with a usage less than 10 pieces a year (see e.g. Ben-Duya, Duffuaa, Raouf, Knezevic & Ait-Kadi (2009, p. 199)). Within SPO, differentiation of forecasting methods could easily be realized, since Xelus also offers techniques like Croston's Interval Smoothing Constant. Croston's method is one of the suggested forecasting methods for slow moving intermittent demand, but SPO currently does not select this method as one of the option Xelus needs to take into account. After selecting the methods, Xelus indicates which of the selected methods should be used as preferred forecasting method. This indication is based on the forecast error, being the smallest deviation between the forecasted usage and the actual curve, measured in terms of the Mean Squared Error (MSE).

For day-to-day forecasting, the applied characteristics are based on the forecasting technique which is selected, but in all cases historical usage is used as input. When the forecast is used in the planning process, different characteristics become important. Stock replenishment orders will be placed based on the forecasts, the actual usage, the EOS date and the current inventory position for new products, products that are available for repair and products that are available for repair under warranty. Order quantities, delivery dates and the source are checked.

For the forecasting process, SPO does not apply performance measures to determine the actual performance of the forecast generated by Xelus. SPO does measure the performance of the actions taken based on the created forecast. The performance measures used are the Parts Availability Level (PAL) and Parts Delivery Time (PDT). PAL represents the percentage of customer orders that can be fulfilled from stock at the moment the order is placed, based on an hierarchical customer order fulfillment strategy, in which a strict location sequence is used to check the availability of the requested FRU. If the FRU is not available at any of the locations in the sequence, this is seen as a PAL loss. PDT is the percentage of orders for which the correct FRUs are delivered at the correct location within the correct time window. For both PAL and PDT target values are determined, which should

be met. These target levels are set for a specific year and differ per division. Based on the results for PAL and PDT, actions might be undertaken to improve the results.

Besides PAL and PDT, SPO also creates a list with all FRUs that are considered to be an emergency (the FRU is not on stock while a customer order is placed) and a list with overdue orders, which are not fulfilled within 2 weeks after the customer order is placed. The FRUs on both lists are monitored carefully and appropriate actions are undertaken to solve the problems as soon as possible.

The entire process is based on the assumption that Xelus automatically selects the most accurate forecasting method from the ones chosen, based on the smallest MSE over the last 2 years, divided in time buckets of 4 weeks. However, Xelus uses fixed parameters in the calculations of the forecasts. For example, when moving average is selected, this will always be a 6-months moving average, although different parameter settings might result in a forecast with a lower MSE. As a result, the selected forecasting method might not be the best possible.

A requirement for the forecasting process is, that manual actions should be minimized due to the large amount of FRUs SPO has in stock. This requirement also relates to the planning process.

3.1.2 New products forecasting

For new products, forecasts are created manually, based on the planner's experience with previous FRUs, the importance of the FRU with respect to the functioning of the machine and data about the expected IB, when available. At the moment all processes for the new products are in place and most errors are solved, the FRU will be forecasted in the normal process. The measurements used to monitor the new products process are the same as those for the day-to-day forecasting. However, new products receive more attention in the daily process.

An important assumption in the new products forecasting process is, that the usage of a new product is similar to previous items, like predecessors of the product or items from other machines that have a similar description or function. A potential threat when applying this assumption can be that the usage pattern is substantially different, due to e.g. an extensive marketing campaign or technology changes creating a substantial increase or decrease in usage.

3.1.3 LTB forecasting

When FRUs are used in manufacturing and/or service, the possibility exists that the supplier stops producing the FRU, and an LTB is necessary. LTBs are done in cooperation with all interested regions worldwide, and occur in 90-95% of all cases before IBM stops using the FRU in the production of new machines. When this is the case, manufacturing is responsible for coordinating the process and stocking the FRUs, and SPO supplies the required information. When production has stopped, SPO is responsible.

In both situations, the PLCM department has to forecast the amount of FRUs SPO expects to use in EMEA during the RSP of the FRU. To determine the total requirement, the forecasted need will be reduced by the on-hand inventory, the amount of FRUs on order and the expected amount of FRUs that can be retrieved from repair, remanufacturing or dismantling.

To be able to forecast the future need, information is required. First, a forecast value for a single month is determined. This forecast can be automatically determined by the program used for the LTB, but is often recalculated by the planning department in Hungary. As a second step, a decline

rate will be determined. Despite the name, the rate can also incline or be stable, depending on the PLC phase the FRU is in at the moment of the LTB. Every division has a standard decline rate, which in most cases is a linear decline rate with a yearly decline of 15 percent. However, this decline rate is not representative for all FRUs. As a result, the decline rate can also be determined based on the expected changes in the IB. The information regarding the expected changes in the IB per machine model is provided by the Service Planning department, based on e.g. information about customer orders and service contracts. By calculating the total IB per year, and setting the IB of the year the LTB is performed equal to 100%, the decline rates are determined by calculating the percentage of the IB still in place related to the IB in the LTB year. Using changes in the IB as indicator for the decline rates implies the assumption that changes in the IB and changes in the demand are correlated with each other.

The amount of FRUs that IBM expects to use during the RSP is represented by the gross LTB need. In the calculation of the gross LTB need, the following information is required:

- The RSP, or number of years until EOS (t)
- The number of months FRU i will be serviced in year t (m_{it})
- The month forecast for FRU i at the moment the LTB calculation is made (f_i)
- The decline rate for FRU i in year t (d_{it}).

When all data is available, the gross LTB need can be calculated by the following formula:

$$Gross\ LTB\ Need_i = \sum_{t=current\ year}^{Year\ of\ EOS} m_{it} * f_i * d_{it} \quad (1)$$

When the gross LTB need is known, the net LTB need can be calculated. This is the amount of FRUs SPO will actually acquire in the LTB. The net LTB need is determined by subtracting the current on-hand inventory, the FRUs on order and the expected repair supply from the gross LTB need. The expected amount of FRUs from repair are forecasted based on data from the Engineering department regarding the return rate of broken FRUs and the amount of FRUs that pass the quality requirements after repair, called the technical repair yield. This can result in a positive or negative net LTB need. As with the day-to-day forecasting, there is no distinction made in the LTB calculation approach for fast and slow moving items. The current approach is suitable for fast moving items, but a standard decline for slow moving items might result in forecasted needs that are inaccurate.

The performance of the LTB forecast is not monitored by specific measures. However, the amount of financial reserves on FRUs that are already EOS or will reach their EOS date within a year, and the number of times an FRU is labeled as a no-source item can be used as an indication. Financial reserves represent the financial savings of an organization, intended to help meet future financial needs the company might encounter, e.g. for scrapping FRUs due to overstocking. When there are financial reserves for a specific FRU, this can be an indication of a large amount of inventory, possibly due to an overestimation of the LTB need. When an FRU is labeled as a no-source item, SPO does not have the FRU in inventory and cannot attain it using other sources. This could be an indication of an underestimation of the LTB need.

For the LTB process, a requirement is that well founded decisions can be made with respect to cost, risk and service. To be able to make a good decision, the best possible forecast with the available

information at that point in time should be created, to minimize the cost over the RSP, while being able to meet the service targets.

3.1.4 The role of forecasting processes in Inventory Management

Cost reduction is important within IBM, and reducing cost related to inventory is one of the ways in which SPO tries to achieve this. To restrict and/or reduce the inventory, inventory targets are created that define the total value that all FRUs in inventory combined are allowed to have. For new products forecasting, the value of the inventory that should be added to be able to service the new product should therefore stay within boundaries. These value boundaries also apply to orders placed on existing FRUs.

Another area in which cost reduction activities are undertaken is in the area of the financial reserves taken to cover future cost of scrapping inventory. To determine the amount of financial reserves SPO should take, a set of rules is determined by accounting. This set of rules specifies which financial reserves need to be taken, and the percentage of the FRU value that the financial reserve should be taken for. There are three types of financial reserves, being:

- Excess: Reserves taken for FRUs that have inventory in stock for more than five or seven years (depending on the division), based on the calculated actual yearly usage, using eighteen months historical usage data
- Surplus: Reserves taken for FRUs that had no usage in the past eighteen months
- EOS: Reserves taken for the on-hand inventory of FRUs that have reached their EOS date

Excess and surplus reserves can change over time, because the FRU is still serviced. EOS reserves cannot change once taken, they can only be used to cover the cost of e.g. scrapping. This makes the reserves an important measure in the LTB process. When an LTB forecast results in a positive need, and there are already high reserves for the concerned FRUs, acquiring parts will only result in additions to the reserves. This can for example occur when there are surplus reserves, but the RSP of the FRU is longer than the amount of years the current inventory can cover, so additional inventory is required. In these situations, it is very important to have forecasts that are as accurate as possible. This creates the opportunity to make well founded decisions, to prevent acquiring FRUs solely for scrapping purposes.

3.2 Commodity Based Lifecycle Forecasting

As mentioned in section 2.1, different improvement initiatives are developed within SPO, all focused on reducing costs. One of the initiatives is CBLF. This approach is based on the assumption that FRUs from a specific commodity within a specific division follow a similar usage pattern, and this specific usage pattern can thus be used in the forecasting process.

The CBLF pilot is currently initiated for the division Lenovo and a selection of FRUs belonging to System X, based on the date of birth of the FRU, the commodity the FRU belongs to and the actual usage per period. PLCs are in most cases created for a specific commodity, implying the assumption that the PLCs of FRUs in a commodity have equal length and usage patterns, resulting in a similar PLC for each FRU in the group. To create the PLC, a number of actions are required.

First, data regarding the date of birth of the FRU, also called the General Announcement (GA) date, and the actual usage per period needs to be retrieved from the databases. As date of birth currently

the introduction date of the first available FRU is selected. During the time the FRU is on the market, it can occur that the FRU is substituted by a newer FRU. This newer FRU is called the successor, and can be introduced e.g. due to changes in technology. The substituted FRU will then become the predecessor of the newly introduced FRU. In case an FRU has predecessors, the date of birth will become the introduction date of the oldest predecessor of all the successors. A graphical example can be found in Figure 3-1, in which the date of birth for the FRU with the normal, dotted and striped usage line are all equal to the date of birth of the original FRU with the normal usage line.

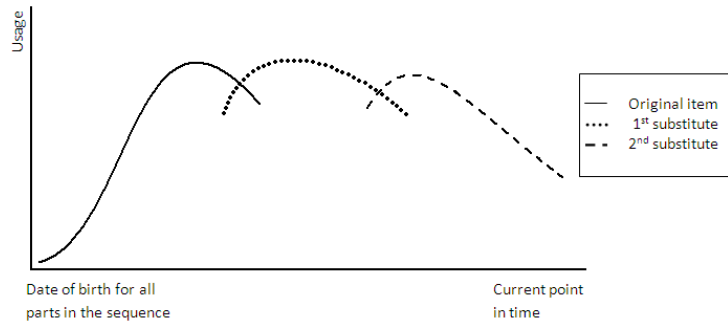


Figure 3-1: Date of birth determination in PLC forecasting

The PLC is determined on commodity level, based on the usage of the FRUs that belong to the commodity. The usage is currently retrieved for four-week time buckets, called periods, within a specific time frame. The time frame for which the data is available ranges from October 2004 until December 2010. A drawback of this approach is that for most FRUs data is only available for a section of the time frame, because introductions can be before or after October 2004 and/or EOS dates can be before or after December 2010. Furthermore, the formal introduction date is not always the point in time at which usage started. As a result, we cannot determine with certainty that the usage pattern we work with is correct, which could result in incorrect forecasts. Newer historical usage data is available, but not yet incorporated into the pilot. This is due to the large amount of manual work required to create the PLCs in the current CBLF model. This is a large drawback, and creating an easier to use model could result in time savings.

Due to the different introduction and EOS dates, the FRUs reach different PLC phases at different points in time. When creating a PLC based on the usage patterns of different FRUs, the phases should be aligned, such that the usage patterns in the different phases are comparable. To realize this, the usage per period is shifted in time, or reallocated, such that the position of the FRU matches the position in the curve it would have reached when full usage data would have been available.

After reallocating the usage such that the PLC phases are aligned, the usage needs to be indexed. The reason for indexing is to make sure that FRUs with high usage per period do not influence the PLC curve more than FRUs with hardly any usage per period. Indexing is done by dividing the actual usage in a period (U_t^a) by the average usage per period for a specific FRU ($\bar{U}$). Due to the fact that there can be no orders in one or more periods in the timeframe, the average usage per period is calculated over the entire timeframe by dividing the sum of the actual usage over the timeframe by the number of periods in which orders could have been placed ($\sum_t P_t^a$). Placed in a formula, the indexed usage per period for an FRU is calculated according to the formula:

$$U_t^i = \frac{U_t^a}{\bar{U}} \quad (2)$$

With:

$$\bar{U} = \frac{\sum_t U_t^a}{\sum_t P_t^a} \quad (3)$$

The rough version of the PLC is established using two different approaches, namely summing and averaging the indexed usage for the FRUs belonging to the selected commodity, and plotting the results over time. Due to the different introduction and EOS dates, the probability that data is available for all FRUs in all periods is very small. As a result, summing the indexed usages can have impact on the accuracy of the PLC, since the sum can differ based on the amount of FRUs for which information is available. Therefore, if additional data is added to the current method, the PLC might change. The differences when adding additional data to average indexed usage will be less, as illustrated in an example in appendix C.

The graph created based on the indexed usage is very capricious. To make the resulting PLC easier to work with, a trend line will be added to smooth out the PLC. Currently, a polynomial trend line is chosen to determine the smooth PLC curve. Often a polynomial of the 6th degree is chosen, which makes the resulting trend line more accurate, but the formula more difficult to work with. An example of the resulting CBLF curve for a hard disk and the corresponding polynomial equation can be found in Figure 3-2.

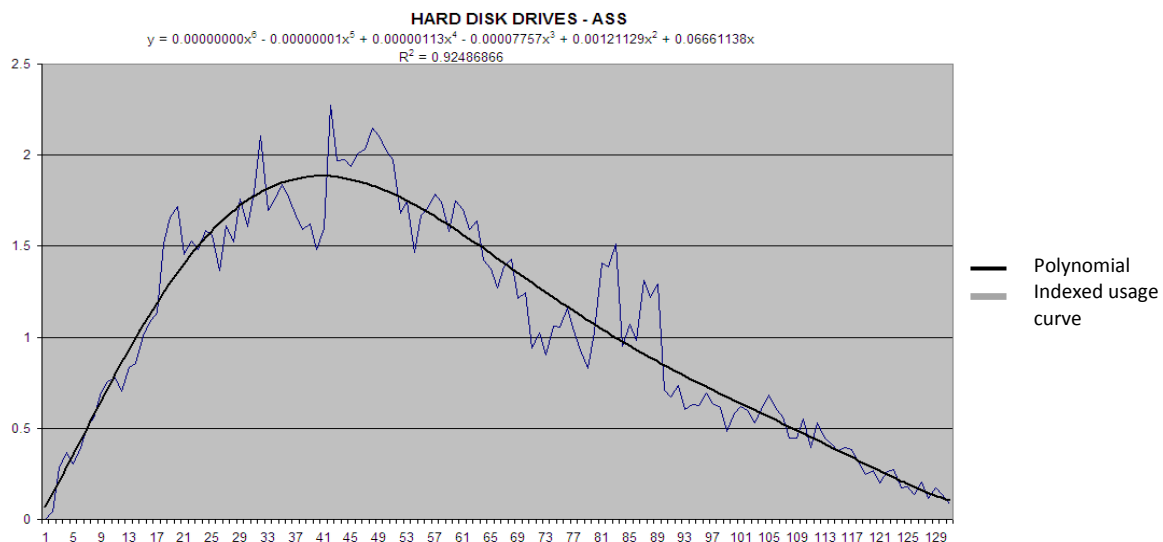


Figure 3-2: Example of the result of CBLF for a hard disk

A large disadvantage of the added polynomial is that the polynomial trend line has a large number of parameters that need to be estimated, which can decrease the accuracy of the formula. Besides that, the polynomial trend line can also have negative values, while negative usage is not possible. Currently this problem is avoided by setting the length of the PLC equal to the time bucket in which the usage becomes negative for the first time. However, this approach implies that PLC forecasting might not always be possible until EOS. A solution for both drawbacks can be to use one of the available standard distributions which cannot have negative values and have less parameters that need to be estimated, e.g. the lognormal, gamma, weibull or beta distribution.

The equation belonging to the polynomial is used in calculating the percentage differences in usage for the desired periods, which in turn are input to forecast the usage in future periods. However, this method will only forecast the mean usage, and does not give any indication of the variability or the forecast error. As a result, it is not possible to give an indication of the safety stock that might be necessary to cover the differences between the forecasted and the actual usage over the RSP.

The determined percentage differences will be used as decline rate in the calculation of the gross LTB need according to equation (1). To calculate the percentage differences, the usage in the period in which the calculation is made is set to be the usage in period 0, the reference period. Then for every future period, the forecasted percentage difference in usage can be calculated with the following formula:

$$PD_t^f = \frac{U_t^f}{U_0^f} * 100\% \quad (4)$$

where PD_t^f is the forecasted percentage difference for period t and U_t^f representing the forecasted usage in period t, based on the formula corresponding to the polynomial of the PLC.

When the percentage differences for all periods are known, they can be multiplied with the usage of the reference period to calculate the forecasted usage in the periods.

The described method is implemented in an Excel work file, which could be used for LTB calculations, and for some FRUs the method is implemented in the planning program Xelus. Both implementations are in the start-up phase. However, the goal is to increase the roll out area to a broader range of commodities and divisions.

For the PLC forecasting, currently no indicators are in place to determine the accuracy of the PLC. As a result, the accuracy and the variability of the usage in the PLC is not calculated.

3.3 Data issues

Before being able to analyze the current performance, a number data issues need to be addressed. One of the issues is that, due to the different data sources that are used, some values are in US Dollars, while others are in Euros. To be able to compare these values, the decision is made to transfer all US Dollar values into Euros, based on the exchange rate at the point in time the value is determined. Using the historical exchange rates, based on Rateq.com (2012), will provide a more accurate indication of the FRU value, due to large changes in exchange rates over time.

Another issue is the issue of missing data. To be able to analyze the FRUs thoroughly in the remaining steps of the research, a large amount of data is required. When vital information, like historical usage data and introduction or EOS dates, is missing and cannot be retrieved via other channels, the FRU will be excluded from the research. Due to the large amount of FRUs and the availability of data for most FRUs that are incorporated in the research, this is not considered to be a problem that will influence the research extensively.

A third issue concerns the historical usage. Due to the registration approach of SPO, in which good returns are seen as negative usage, negative usage can be registered in the systems. Good returns are FRUs that are ordered by the Customer Engineer because they could have been necessary to repair the machine, but turned out not to be the source of the problem. As a result, these FRUs are not used and are returned to SPO, where they will be stocked until a new order comes in. Based on the definition of usage being the actual amount of FRUs used to keep customer machines up and running, good returns are not part of the usage. Therefore, the negative values will be subtracted from the actual usage.

3.4 Case selection

Due to the large amount of FRUs SPO services, the research will only be executed for a selection of the FRUs. In this section we will first describe the procedure which resulted in the selection of the FRUs, followed by an overview of the characteristics and usage patterns covered by the selection.

3.4.1 Case selection procedure

The main research question focuses on forecast accuracy. As a result, performance measures will be determined to be able to answer the research question. The current performance will not be determined for all FRUs SPO has in inventory, to limit the scope and to make sure performance will only be determined for the areas that are important for the research. Therefore, the following selection procedure is created:

1. Exclude divisions that will leave IBM or divisions for which no LTBs are executed.
The division Lenovo will be transferred to another party by the end of 2012. As a result, this division will be excluded from the research, since SPO will not be able to benefit from the possible improvements. Besides Lenovo, the division MVS will be excluded. For this division, SPO does not perform any LTBs, this is the responsibility of the vendor of the part.
2. Exclude the FRUs for which an LTB calculation is not executed and/or data is unavailable
To be able to test the effects of the research outcomes on LTBs, we will only focus on FRUs for which an LTB is already executed. Information can be gathered over FRUs for which an LTB is executed in the last 10 till 15 years from the databases. However, detailed information over the historical LTBs is not always available, making an additional selection based on the available information necessary.

Since detailed information about the actual usage is available from August 2007 to April 2012, the selection will contain FRUs that have an LTB finished in August 2007 at the earliest and March 2012 latest. In the selection process, only LTBs that are completed will be taken into account, to prevent the possibility of changing needs during the research. A second selection criterion is that all required information needs to be available. Both criteria are required to be able to do a detailed analysis of the LTB decision.

LTBs can be executed for different reasons, e.g. because an external supplier announces a production stop, but a recalculation of inventory levels by internal manufacturing departments might also be a reason. This creates the possibility that the same FRU is included in the selection more than once. Although a large part of the characteristics for these FRUs will be the same in both LTB calculations, the decline rate and the forecast can differ, because they depend on the period in time and the position in the PLC. All LTBs in which the FRUs are involved will be incorporated into the selection, due to the differences and since this also creates the opportunity to assess whether there are performance differences between the LTBs at different points in the PLC.

Before the selection procedure is applied, approximately 260.000 FRUs are categorized as potential FRUs to be incorporated into the research. After the removal of the divisions MVS and Lenovo, approximately 110.000 FRUs remain. Of these 110.000 FRUs, an LTB is executed for approximately 7.300 FRUs. When excluding the FRUs for which the LTB does not meet the requirements, 2419 FRUs remain in the selection, divided over 550 LTBs. Excluded LTBs are LTBs that are cancelled, LTBs that regard a component of an FRU instead of a complete FRU, LTBs that are not completed yet and sizings. In a sizing, a rough estimate is made regarding the amount of FRUs necessary to cover the

need until EOS, without acquiring the result. An example of a reason for a sizing is when the manufacturing department has a large quantity of an FRU in stock, and wants to dispose a part of their stock. Then SPO is asked to give an indication of the amount of FRUs they expect to use until EOS, which will be taken into consideration by the manufacturing department in their scrap decision.

The selection of 2419 FRUs will be used in the remainder of the research, and will be divided in a selection of approximately 80% for model building and 20% for model testing. The reason for the division is that if FRUs are used for both model building and model testing, the model will be created such that it will fit the data in the best possible way. As a result, it will not be possible to assess the performance in a statistically independent manner. Therefore, the model building selection will be used in the determination of the current performance and in the creation of an improved forecasting model. The test selection will be used to test which potential improvements can be realized by applying the developed model.

Since categorization on division is the leading approach within SPO, it is considered to be valuable to be able to test the suitability of the created model for all divisions in the selection. To be able to test the created model for all divisions in the selection, the selection of 2419 FRUs is ordered by division, and 483 FRUs are transferred to the test selection based on randomly generated numbers. This results in a model building selection of 1936 FRUs.

As a result of the sorting, the probability that all FRUs of a specific division are incorporated in only one of the two selections is minimized. By applying random selection, the probability that the characteristics of the selected FRUs for model testing are similar to those of the selected FRUs for model creation is high. An advantage of similarity of the selections is the possibility to extrapolate the results of the tests to FRUs that have similar characteristics as the selection, like future similar FRUs.

3.4.2 General overview of the model creation selection

Before being able to analyze the model building selection and draw conclusions based on the results, we first investigate the 1936 FRUs that are assigned to the model building selection, to have some background information regarding the selection we will focus our research on.

The 1936 FRUs are combined in 495 different LTB calculations, which had to be executed while the FRU was still used in the production of new machines in 73,3% of all the cases. With respect to the net LTB need, the conclusion that SPO does not need additional inventory is reached in 69,7% of the cases, indicating that overstocking already occurs in the time before the LTB calculation is executed. This is further confirmed by the determination of the minimum net LTB need, which is -4928 pieces, meaning an overstock of 4928 FRUs on top of the gross LTB need in the RSP. Furthermore, 1694 FRUs are labeled as FRUs with financial reserves. For 1214 FRUs, scrap actions are executed after the LTB is executed. These results indicate that on average the amount of pieces acquired in the time an FRU is on the market is too high.

With respect to the time between the introduction of the FRU and the moment an LTB is executed, a large range can be found in the selection. For some FRUs, the LTB is already executed in the year in which they are introduced, while for others there are 24 years between the introduction and the LTB. When we look at the RSP, the time boundaries regard a period of less than a year up to 16 years to EOS. These ranges imply that LTBs are executed in different PLC stages and for different planning

horizons. This introduces the desire that the created model is able to forecast the PLC in different PLC stages for different planning horizons, without large adaptations.

From a cost perspective, the value of the FRU and the procurement costs for the LTB are assessed. The value of the FRU is based on weighted average costs. Within the selection, FRU values range from €0,01 till €38.599. Regarding the LTB procurement cost, the minimum LTB procurement value is €0. This corresponds to the FRUs for which SPO already has more in stock than they expect to use until EOS. The maximum value of the LTB procurements is slightly over €7M.

An overview of the minimum, maximum and average value regarding the age at the moment of the LTB, the RSP, the value of the FRU, the net requirement and the LTB procurement costs can be found in Table 3-1.

	Minimum	Average	Maximum
Age FRU at the moment of the LTB	<1 year	5,3 years	24 years
Remaining service period	<1 year	6,1 years	16 years
Value FRU	€0,01	≈ €948	≈ €38.599
Net requirement	-4.928 parts	≈ 44 parts	11.323 parts
LTB procurement cost	€0	≈ €27.500	≈ €7M

Table 3-1: General range overview of the initial selection

With respect to the LTB process, for 789 FRUs a standard decline rate is used. For 19 FRUs, the decline rate first increases, and for 24 FRUs the decline rate is stable. The remaining 1104 FRUs have a decline rate that declines based on the expected IB developments.

3.4.3 Identifiable partial usage patterns

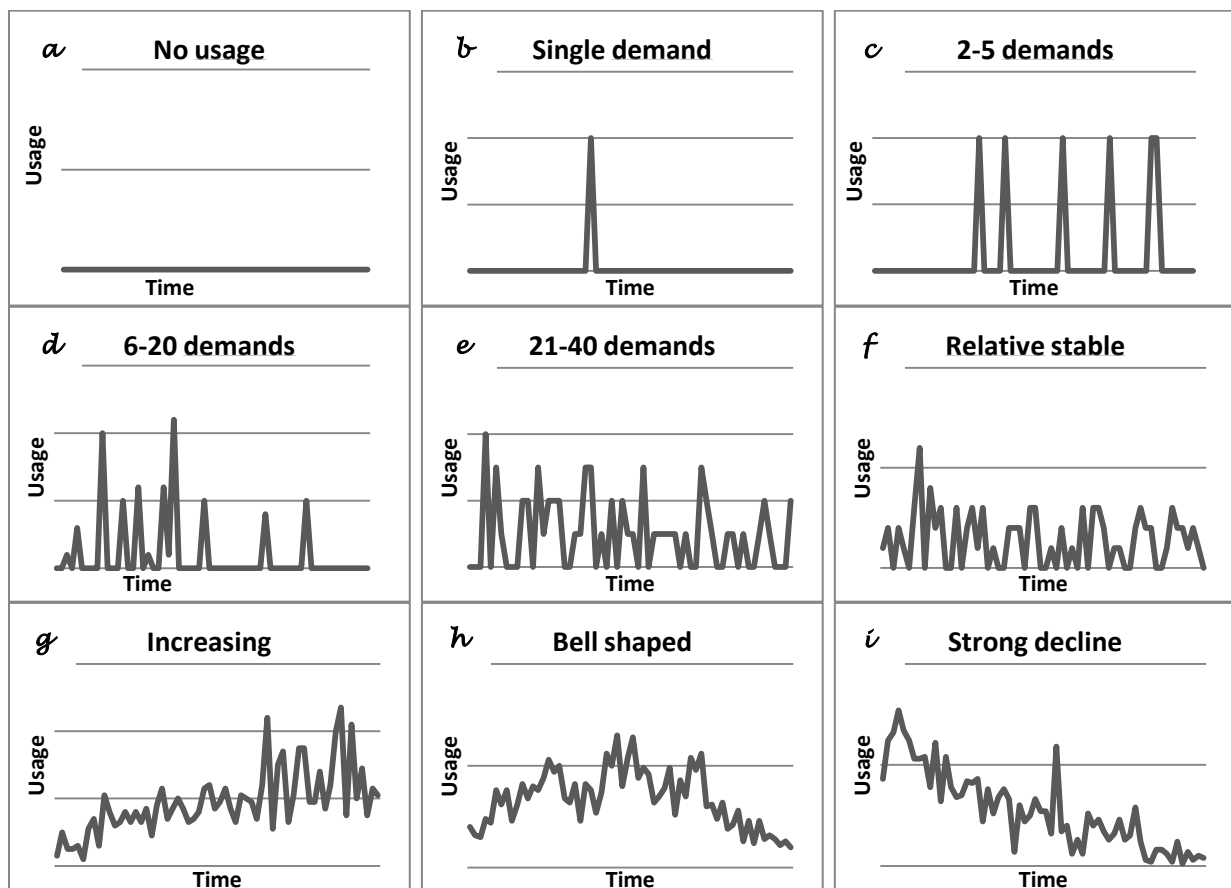
In our research we focus on clustering FRUs such that the PLC of the FRUs in a group is similar, to simplify forecasting processes. Different options are considered in literature regarding existing usage patterns and how to cluster them. In her article, Wood (1990) describes 15 different usage patterns. Meixell and Wu (2001) state that products follow a few identifiable usage patterns, and they can be clustered according to these patterns. In another article, Wu et al. (2006) discuss a clustering approach based on mean shipment quantities, shipment frequencies, volatility or skewness of usage patterns.

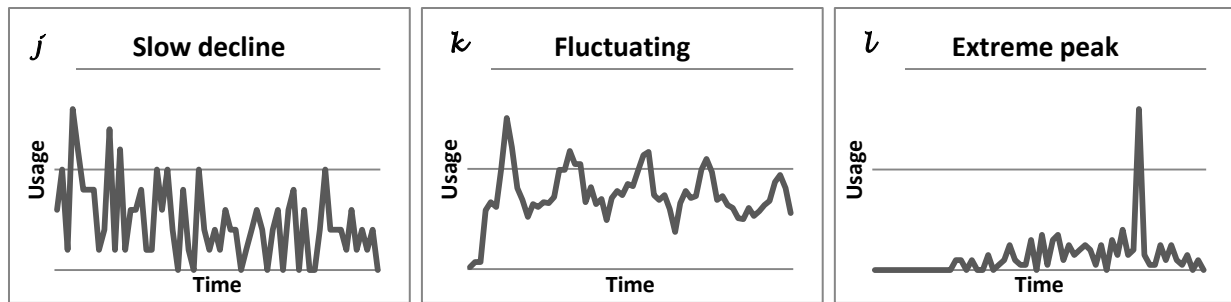
Different usage patterns might need different forecasting methods, making it important to know which usage patterns can be found in the model building selection. The decision is made to investigate whether there is also a limited set of usage patterns to be identified within the product set of SPO. Therefore the partial usage patterns over the period ranging from August 2007 till April 2012 are created. Since this is only a limited time horizon, complete usage patterns are not likely to be found, due to the fact that most FRUs have an expected time on the market that exceeds 5 years and because there is uncertainty regarding the actual starting point of the usage. However, partial usage patterns can be used as an indication of the PLC stage the FRU might be in and whether solely the standard PLC pattern is present or some of the other usage patterns as well.

In the process, the point in time at which the LTB is executed is not taken into account, because this is not relevant for the identification of the usage patterns that could be available. However, when complete usage patterns can be created, the point in time at which the LTB is executed could provide a valuable indicator of the future usage.

To limit the amount of usage patterns, the partial usage patterns are combined into 12 different groups of partial usage patterns based on graphical resemblance. For every partial usage pattern, an example can be found in Figure 3-3. Since this is an example of an FRU corresponding to the pattern, not all FRUs in the group will have the same usage per period, but the underlying pattern will be similar. Due to the large amount of partial usage patterns in which usage is not continuous, the decision is made to divide these partial usage patterns over a limited set of groups. Using Syntetos, Boylan & Croston (2006), who mention that intermittent PLC patterns can be characterized based on the amount of periods in which usage occurs, the division is made based on the amount of periods with usage in the 5 year horizon. The boundaries are set such, that it should be possible to roughly categorize the partial usage patterns based on the observed usage realization, without having to determine the amount of periods in which usage occurs. As a result, the range is small for the case in which usage occurs seldom, and increases as the number of periods with usage increases. This differentiation is created to be able to test if there are characteristics indicating a certain range of usage occasions, and which forecasting method is best suited for intermittent usage patterns. In appendix D an explanation of the partial usage patterns and the reason for appearance can be found.

When we try to assess which usage patterns might be present, the traditional PLC pattern is an important option to consider. This is, because the traditional PLC shape can be created out of a combination of the bell shaped and the strong and/or slow decline pattern. A pattern with slow increase and slow decrease could also be present, based on the increase and slow decline pattern. The fluctuating, stable and intermittent usage patterns are other options to take into account. As a result, the large amount of combinations that can be created based on the partial usage patterns make it impossible to actually determine which usage patterns can be found at SPO.



Figure 3-3 j - l Identified usage patterns

The 12 identified partial usage patterns from Figure 3-3 are not evenly distributed over the 1936 FRUs, as can be found in Table 3-2.

PLC pattern	% occurrence in the model selection
No usage	23%
Single demand	9%
2-5 demands	12%
6-20 demands	19%
21-40 demands	10%
Relatively stable	3%
Increasing	2%
Bell shaped	5%
Strong decline	4%
Slow decline	8%
Fluctuating	4%
Extreme peak	1%

Table 3-2: Percentage occurrence usage patterns

In the selection, partial usage patterns with no or incidental usage are most common. These patterns could occur at the end of the PLC where the usage is expected to be low, but also at earlier points in the PLC. When this is the case, the standard PLC pattern might not be appropriate for these FRUs. The declining patterns, which one could expect to be most common due to the fact the LTB is expected to occur near the end of the PLC, only represent 12% of the selection. However, this percentage could be higher, when LTBs for the bell shaped pattern occur after the peak.

3.4.4 Conclusion

We have created a model building selection consisting of 1936 FRUs and a model testing selection of 483 FRUs. The model building selection will be used to assess the current forecasting performance and covers a large range of characteristics values regarding the time between introduction and LTB, the RSP, FRU value, net requirement and LTB procurement cost. Furthermore, we have identified 12 different partial usage patterns, indicating that different usage patterns might be present in the selection. We cannot establish the usage patterns actually present within SPO, due to the large amount of possibilities and the limitations posed by the limited amount of historical usage data.

3.5 Current performance

As a starting point of the research, the current performance will be determined. Knowing the current performance will provide the opportunity to determine in which areas improvements could be realized. First, the overall performance of the FRUs in the model building selection will be determined, followed by a performance breakdown for the identified partial usage patterns and the commodities.

3.5.1 Performance general selection

To get an indication of the performance of the LTB calculations of SPO, the performance for the LTBs in the model building selection will be determined in terms of the forecast accuracy. Forecast accuracy resembles the extent to which the forecast accurately predicts the usage in a period. The higher the forecast accuracy, the better the usage is predicted. Since the forecast and the actual usage data are known, the model fit will be determined, which represents the extent to which the created forecast model actually fitted the realized usage curve.

Different techniques are used to determine the model fit. First, the bias of the forecast is measured, which gives the average error size and can be used as an indication to assess whether SPO generally over or under forecasts the usage in the RSP. Second, the Mean Absolute Percentage Error (MAPE) is determined, which gives an indication of the average size of the forecast error in percentages of the actual usage. Using this measure provides the opportunity to assess whether the forecast errors are similar for FRUs with different usage values. A drawback of this measure is, that it will give infinite values when the actual usage is zero. Finally, the Mean Absolute Deviation (MAD) is determined. The MAD gives an indication of the size of the forecast error. The formulas used to calculate the different measures are:

$$Bias = \frac{1}{n} \sum_{t=1}^n (x_t - \hat{x}_t) \quad (5)$$

$$MAPE = \frac{1}{n} \sum_{t=1}^n \frac{|x_t - \hat{x}_t|}{x_t} * 100\% \quad (6)$$

$$MAD = \frac{1}{n} \sum_{t=1}^n |x_t - \hat{x}_t| \quad (7)$$

with n being the number of periods over which the model fit is calculated, x_t being the actual observed usage in period t and $\hat{x}_t$ resembles the forecasted usage in period t , determined at the moment the LTB calculation is made.

Using equations (5)-(7), the forecast accuracy is computed for the 1936 FRUs that are incorporated in the model building selection. Based on the bias, 1051 FRUs show a negative result, indicating that the forecasted value is on average higher than the actual usage, with an average difference of 1,41 pieces per 4-week planning period. For 528 FRUs, the bias is positive, indicating that on average the forecast underestimates the realized usage. Here the amount of pieces underestimated is on average 2,59 per 4-week planning period. The remaining 390 FRUs have a bias of zero. This is the result of the fact that for all these FRUs, the forecast was zero usage over the RSP, which has also been realized. With a forecast of zero and a usage of zero, the bias will be zero as well.

The MAPE is calculated to be able to compare the deviation between the forecasted and actual usage. For FRUs that have no actual usage over the period since the LTB has occurred, MAPE will give infinity as result. For these FRUs an indication of the forecast error will be given by the MAD. Based on a guideline for performance classification given by Chien, Chen & Peng (2010), as can be found in Table 3-3, an indication will be given to be able to categorize the performance based on the results of the MAPE calculations.

MAPE performance range	Performance classification
< 10%	Best forecast
10-20%	Good forecast
20-50%	Reasonable forecast
> 50%	Inaccurate forecast

Table 3-3: MAPE performance classification according to Chien et al. (2010)

For comparability reasons, the average MAPE represents the average percentage difference between the forecast and the actual usage per 4-week period. This approach will cancel out the possible effects of a different amount of periods used in the calculation, which is important because the LTB can be executed at any point in the 5 year time horizon. This creates the possibility that the performance is determined over e.g. a 4 year time horizon or a 1 year time horizon, which could severely influence the results. For the average MAD, the same approach is taken.

For the 1309 FRUs that have usage in the period since the LTB, the average MAPE value is 32% per 4-week planning period. So on average, the forecast can be classified as reasonable. However, the maximum MAPE value is 858% and 282 FRUs have a MAPE value equal to or over 50%, indicating an inaccurate forecast. This clearly indicates that there is room for improvement within the current forecasting procedure.

When looking at the MAD values, 390 of the 627 FRUs having no usage since the moment of the LTB have an MAD value of 0. These 390 FRUs are the FRUs for which the forecast was zero and the actual usage turned out to be zero as well. 7 FRUs have an MAD value of more than 1, the remaining FRUs have an MAD value between 0 and 1. This leads to the conclusion that in most situations where there is no usage, the created forecast is on average not much higher than the actual usage of zero.

The results of the forecasting accuracy computation for the model building selection points in the direction that in most cases the forecast overestimates the actual usage. It is assumed these results can also be found outside of the model building selection. Concluding from these results, it can be stated that the overall forecasting accuracy provides room for improvement.

3.5.2 Performance breakdown per partial usage pattern

Besides looking solely at the overall performance, we also determined the performance per partial usage pattern as identified in section 3.4.3, to see whether certain partial usage patterns are forecasted more accurately than others. To do this, the bias, MAPE and MAD will be used, as was the case for the general performance assessment.

With respect to the difference between the forecasted and actual usage, the largest MAPE values can be found in the fluctuating, increasing, bell shaped, slow and strong decline patterns, as displayed in Figure 3-4. The high deviation for both declining usage patterns is surprising, since in most LTBs a decline rate is used that solely decreases. As a result, the MAPE value gives rise to the question whether the current decline rates actually resemble the decline in the usage. The large MAPE values in the increasing and bell shaped usage patterns are expected, since the SPO team indicates that it is difficult to forecast increasing usage patterns. One of the reasons for this is that decline rates with an increasing section are rare, most likely causing large deviations for usage patterns with an inclining section. The partial usage patterns with low usage appear to be forecasted quite accurate, which is a positive result, since they represent a large percentage of the observed partial usage patterns in the selection.

PLC pattern	# FRU	total value LTB
No usage	444	€65,8K
Single demand	168	€265,3K
2-5 demands	226	€101,6K
6-20 demands	372	€5,3M
21-40 demands	185	€2,4M
Relatively stable	57	€667,7K
Increasing	49	€5,1M
Bell shaped	93	€3,8M
Strong decline	84	€3,9M
Slow decline	146	€7,8M
Fluctuating	80	€6,7M
Extreme peak	29	€150,1K

Table 3-4: Information partial usage patterns

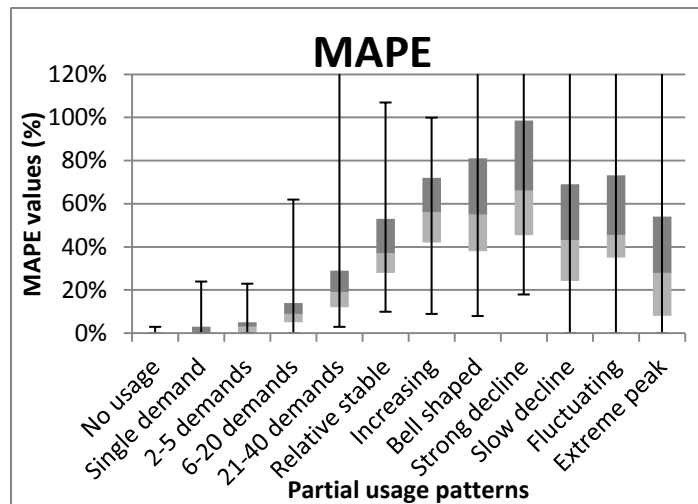


Figure 3-4: Box plot MAPE values per partial usage pattern

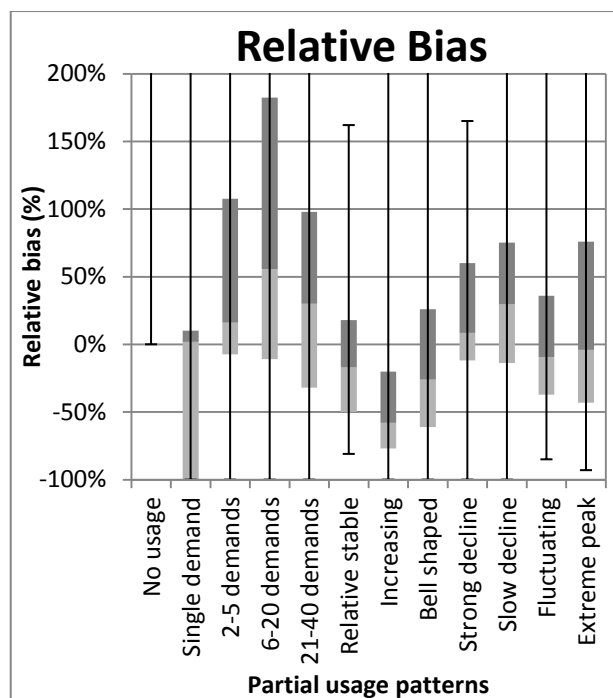


Figure 3-5: Box plot relative bias values

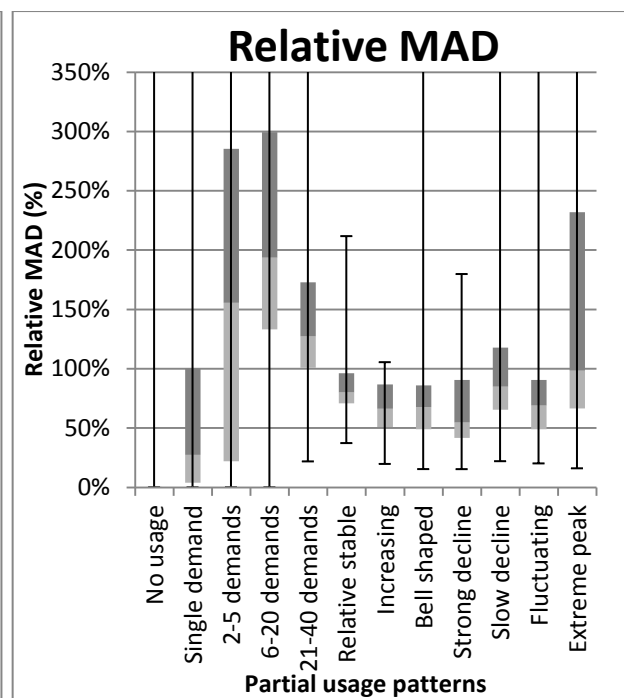


Figure 3-6: Box plot relative MAD values

Based on the performance calculations for the partial usage patterns, we can conclude that the bias ranges between -1,5 and 1,5 for all partial usage patterns, with the exception of the increasing and bell shaped pattern, which have a bias larger than 4. On an aggregated level, this appears to be low. A possible cause is the fluctuation in the usage, which results in positive and negative differences cancelling each other out. However, it is difficult to interpret the impact of the bias on the partial usage patterns. Therefore we calculated the relative bias, which is the bias divided by the mean usage, to get more insight in the impact of the bias on the partial usage patterns. The relative bias values, in which a positive percentage indicates over forecasting and a negative value under forecasting, are represented in Figure 3-5. All partial usage patterns, except the increasing usage pattern, cover a range which includes 0% bias, representing an accurate forecast is possible. It is also clear that both over and under forecasting occurs, but the magnitude is larger in case of over forecasting.

Based on the MAPE values we concluded that the results for the low usage patterns are quite accurate, but the relative MAD values might indicate otherwise. We calculated the relative MAD values by dividing the MAD value by the mean usage, and plotted the ranges in Figure 3-6. From this figure we can see that the percentage deviation is the largest for the low usage patterns, while a value close to zero is desirable. As a result, based on the MAD value we conclude that there is potential to improve the forecast for the low usage patterns, but overall different partial usage patterns are considered to be more interesting based on the MAPE values.

Overall we can conclude that several partial usage patterns have a large probability of having a MAPE value larger than 50 percent, indicating an inaccurate forecast and an opportunity for improvement. Parts with low or relatively stable usage are forecasted quite accurate, with the exception of partial usage pattern 6-20 demand occasions. The low usage patterns are considered to be less interesting due to the relatively low investment value, as can be seen in Table 3-4, and will therefore receive less attention.

3.5.3 Performance breakdown per commodity

In the previous sections, the current performance of the entire selection and the performance per partial usage pattern are discussed. To be able to compare the performance of the current method with the performance that might have been realized with CBLF, we will also assess the performance per commodity, based on the MAD, MAPE and bias.

The FRUs SPO services are categorized per commodity. A commodity is a group of FRUs that have a similar function and might be interchangeable. An example of a commodity are batteries, which are all used to provide power, but exist in different types. Due to the amount of available commodities, a selection is made to limit the number of commodities. In Table 3-5 the commodities for which the total value of the executed LTBs in the model building selection exceeds €1 million are listed. A short explanation of the commodities can be found in appendix E.

Based on the results, as listed in Table 3-5, the performance of the commodity HDD shows the highest deviation from the actual usage per 4-week period. Since HDDs are considered to be a difficult commodity to forecast by the SPO team, due to large amounts of substitutions, combined breakdowns and matrix structures which indicate that specific HDDs are interchangeable, this result is not considered to be exceptional by the SPO team. The commodities CPU and MISC show relatively low values for both MAPE and MAD, leading to the conclusion that they are forecasted relatively accurate. Especially for the commodity CPU this is considered to be a positive result, due to the high total LTB value concerned. For all commodities, the bias is low. This is due to the erratic usage patterns, which results in a bias that is cancelled out by over and under forecasting.

FRU Commodity	# FRU	Avg. bias	Avg. MAPE	Avg. MAD	Total value LTB
CARD	289	0,28	26%	1,64	€4,5M
CPU	284	0,01	15%	0,56	€8,7M
HDD	152	0,02	48%	6,08	€5,5M
MEM	128	1,20	37%	3,71	€7,7M
MISC	298	0,01	10%	0,62	€1,7M
PLNR	86	1,08	28%	4,67	€1,4M
POWR	75	1,02	46%	3,28	€2,5M
TAPE	75	0,75	31%	3,33	€3,0M

Table 3-5: Current performance standard approach, per commodity

Although all average MAPE values are below 50%, the accuracy can still be improved. This could lead to high cost savings, since the selected commodities represent a large section of the total amount invested.

3.5.4 Conclusion

In this section we determined the forecast performance for the model building selection based on MAPE, MAD and bias results. In general, a large section of the FRUs are over forecasted, and MAPE values indicate that the forecasts are inaccurate. Similar results can be found on partial usage pattern and commodity level, although the bias is relatively low for both groups, due to the erratic usage patterns. Concluding we can state that there is a high improvement potential for general forecasts, but also for forecasts specifically designed for a specific subsection of FRUs.

3.6 Improvement possibilities

Based on the analysis of the current way of working, different improvement possibilities can be distinguished. These improvement possibilities are divided into possibilities related to the current way of working and possibilities related to CBLF.

3.6.1 Improvement possibilities related to the current way of working

Currently, LTB calculations are executed by different employees for different divisions. As a result, a difference in approach can be found with respect to the chosen decline rates and the amount of risk or safety that is built in for the specific LTB. This is mainly due to the fact that there is no standard procedure. Besides differences in approach, there are also differences in the way in which the team documents the actions executed. This will make it more difficult to assess whether specific actions are mostly correct or incorrect. It could e.g. be the case that one of the employees always adds an amount of 0,5 FRU to the month forecast used in the LTB. When assessing the forecasting accuracy at a later point in time, without information about the added safety factor, it is impossible to determine whether the addition of the safety factor resulted in a better forecast than would have been the case when the safety factor was not added. The same holds for taking additional risks. Knowing this information can, after assessing the performance over time, result in better guidelines of when to take which actions.

The LTB monitoring system, which is currently being developed, will give the opportunity to assess whether the forecast adaptations, chosen decline rates and expected supply plans lead to an adequate representation of reality. Based on the information about the accuracy of the decisions made at the moment of the LTB over time, a more standardized approach could be developed for specific groups of FRUs or for specific situations. The information could also provide guidelines for situations in which it is more appropriate to take risks or build in extra safety.

Another improvement option is the application of different forecasting methods for slow and fast moving items. This will provide the opportunity to get better insights in the usage patterns for both types of FRUs, and this will help in the creation of more accurate forecasts. It would also provide the opportunity to determine the volatility in the usage patterns, which could be used to determine a range of values in which the actual usage will end up with a predetermined percentage of certainty.

3.6.2 Improvement possibilities related to CBLF

In section 3.5.2 we concluded that the largest potential for improvement lies in several partial usage patterns with high MAPE values. Here we will discuss options in which CBLF can be adapted to improve the forecast.

As described in section 3.2, an improvement could be to always use the average indexed usage instead of the summed indexed usage when the CBLF curve is created. This is expected to lead to more accurate forecasts in case of missing data, as quantified in an example in appendix C. This approach is tested in section 6.3.3.

With respect to the creation of the smoothed PLC curve, it might be possible to use a different function to represent the PLC curve for a specific commodity. This can for example be a lower class polynomial, which reduces the amount of parameters in the equation. Another option could be the use of a standard distribution, like e.g. the lognormal or Weibull distribution, which would eliminate the possibility of having a curve which contains negative usage. The standard distributions might also be used to forecast the usage curve based on actual usage for a specific FRU. This will increase the opportunity to adapt the usage curve to a specific FRU. We will further examine the impact of this approach in section 6.3.1.

In the current approach, it is not possible to adapt the length of the curve. This will not lead to problems if all FRUs in the commodity will be approximately the same time on the market. However, when there is a substantial difference between the time between GA and EOS for a specific FRU and the expected time the FRU is on the market based on the CBLF curve, the points in time in which the FRU goes from one PLC stage to another might not be aligned correctly. Making the curve adjustable will overcome this problem, for example by using percentages of the total usage consumed in a specific part of the PLC. Further information regarding this approach can be found in section 6.3.1.

A final improvement focuses on decreasing the complexity and the amount of manual work necessary to perform an LTB based on CBLF using the worksheet. This could be realized by incorporating macros that will automate the process. This will also reduce the amount of errors that could occur. Performance measures should also be included, to determine whether CBLF leads to improved results. This is important, since there is a probability that not all FRUs follow the described PLC curve, based on the partial usage patterns discovered in section 3.4.3. When information about the performance is known, guidelines can be created regarding the use of CBLF or other methods.

3.7 Conclusion

Within IBM, different forecasting methods are applied in the LTB situation, which are all based on standardizing the expected usage pattern over the RSP. For some FRUs, these approaches yield relatively good forecasting results, but for others the methods provide room for improvement. However, in general the forecasted amounts over the lifetimes of the FRUs are too high.

Besides assessing the performance of the forecasted methods, we also identified 12 different partial usage patterns in the model building selection. Although we were not able to combine these partial usage patterns into complete usage patterns, the results indicate that there is a possibility that FRUs follow more than one type of usage pattern.

4 Model prerequisites and literature background

Before continuing with the research, the prerequisites for the model and a literature framework regarding PLC forecasting are provided. This framework consists of basic information regarding the concept of PLC forecasting, available methods and applied characteristics. It can be used to get acquainted with the topic and in determining what possible methods could be applied at SPO.

4.1 What are the prerequisites for applicability to the situation of SPO?

IBM operates in the IT branch, which is characterized by rapid technology changes and a large diversity in PLCs. Due to the market characteristics, machines or machine components can become obsolete shortly after introduction, as is the case with e.g. laptops. On the other hand, customers also acquire large and sophisticated machines, like mainframes, that should be operational for more than 10 years. Due to the high diversity in the length of the PLC, it would be considered as a pre if the model is able to cover PLCs with different lengths.

With respect to the actual usage of spare parts, large difference can be distinguished. A section of the spare parts is ordered just a few times per year. Another section has a very high usage. In both cases, the usage patterns can be stable or very capricious, with high usage in one period and very low or no usage in others. Therefore, a prerequisite is that the model is either fitted for a specific type of usage pattern or can identify a selection of the most common usage patterns.

Since the resulting model should be incorporated into the current way of working, the computations should be feasible in general applications, like Excel. All models for which more sophisticated (mathematical) programs are necessary will be discarded, because of possible difficulties when the model is embedded in the current way of working.

With respect to the required input for the computations, data should only regard FRU related information, information of similar FRUs and limited additional internal information. Competitor information is not incorporated, because SPO does not have direct access to this information and rapid changes in the industry will result in a high probability of outdated information. Another requirement regarding information is that the model should work with information covering a part of the PLC, since five years of history for the actual usage is available and only two years of historical information for other data. Flexibility in the time period for which a forecast is made could also be valuable for the SPO team, since this would provide the opportunity to derive an aggregated or more detailed forecast, depending on what is required.

Finally, the basic idea behind the computations should be understandable for the SPO team. When the model is understood and accepted, the probability it will actually be applied in the decision making process increases. An easy to use model also contributes to model acceptance, which results in the condition that the amount of steps in the model should be simple, straightforward and limited.

4.2 The basic PLC forecasting concept

The PLC concept started to appear in literature in the 1950s and 1960s, introduced as a marketing tool which could be used to forecast the future requirements and develop appropriate strategies to manage future usage. Originally, the PLC concept is described as a four stage model. The stages are respectively introduction, growth, maturity and decline (Rink & Swan, 1979). These stages form a classical S-shape, as can be found in Figure 2-1. Since the introduction, different authors have come up with other patterns based on their research. Wood (1990, p. 149) gives an overview of 14

patterns besides the classical pattern. These different patterns include e.g. multiple peaks and constantly increasing usage. Meenaghan and O'Sullivan (1986, p. 89) state that the shape and length of the PLC is determined by four clusters of influencing factors. These four clusters are:

- Product characteristics: features of the product itself
- Employed marketing strategies: powerful companies in an industry have more influence on the shape of the PLC than companies with less power
- External environmental factors: factors that influence the shape and length of the PLC that are uncontrollable by the company
- Market-related factors: adoption and diffusion patterns in the market.

Over time, contradicting views on the applicability of the PLC concept emerged. Wood (1990) lists a number of contradicting results of empirical studies, both confirming and disproving the PLC concept. As an overall conclusion, she states that market and technology developments since the introduction of the PLC concept have made the environment more dynamic and complex than was the case during the introduction of the PLC concept, especially in the area of high-tech products and consumer durables, making the PLC concept not the most flexible and scientific approach anymore.

However, since Wood published her article, numerous authors have published articles about PLCs, e.g. Meixell & Wu (2001), Karmeshu & Goswami (2001) and Huang & Tzeng (2008). These articles include reviews and empirical studies that provide evidence of the applicability of the method, but also newly developed methods to create PLCs for different markets. These articles create the impression that the PLC concept is not yet outdated, but might still be able to provide valuable techniques and insights in usage patterns.

4.3 What literature is available regarding PLC clustering?

In their proposed PLC forecasting method, Meixell and Wu (2001) use historical usage for a specified number of periods as a guidance for clustering products. These clusters should be created such that the parts within a cluster are more similar to each other than to parts in different clusters. A similar approach is described in Wu et al. (2006), who mention options for categorization based on mean shipment quantity, shipment frequency, the volatility of the usage and the skewness. Another clustering option is proposed by Syntetos et al. (2006), who state that if the objective of categorization the selection of the most appropriate estimation procedure is, comparison should be done based on an accuracy measure.

Williams (as cited in (Syntetos, Boylan, & Croston, 2006)) created a categorization scheme based on the number of lead times between two consecutive usage occasions and usage fluctuation. With these aspects, four different usage patterns can be distinguished. Eaves (2002, p. 127) revised the classification scheme of Williams into a scheme with five possible types of usage patterns, based on transaction variability, order size variability and lead-time variability.

Besides clustering on usage patterns, different options are applied in literature. Kirshners and Sukov (2008) cluster parts based on the maximum and minimum normalized usage. These maximum and minimum values are placed within a grid, which is divided into clusters. In a study focused on creating clusters of pharmaceuticals that follow similar PLC patterns, Bauer and Fischer (2000) cluster the pharmaceuticals under study based on the estimated regression coefficient. These represent the

parameters of the function that is used to model the PLC, and as such form an abstract representation of the PLC curve.

4.4 Which PLC forecasting methods are described in literature?

In literature, different approaches are available to forecast the PLC of a product. PLC forecasting models can e.g. be based on historical usage data. This can be either data from the product itself, or from previous generations (Lapide, 2002). When data of previous generations is used, reliable results can only be achieved when similarity exists between the products of the different generations. In their article regarding technology adoption in high-tech markets, which includes the industry IBM operates in, Meade and Rabelo (2004, p. 673) state that “industries tend to be relatively stable across product generations”. An example of a method based on data of previous generation is the method of Huang and Tzeng (2008), based on fuzzy regression.

One of the available groups of methods is the group of growth or diffusion models. Growth models are based on the Bass model, which is developed in 1969. This model forecasts the cumulative number of adapters at time t based on three parameters representing the influence of internal and external factors. Different variations on the Bass model are developed, e.g. by Karmeshu and Goswami (2001), who developed an approach that more accurately predicts usage patterns with multiple peaks, or by Chien, Chen and Peng (2010), who introduced a diffusion model which incorporates seasonal effects. Besides the growth models, statistical distributions are also applied to predict the usage patterns, as described in Petrescu (2009).

The basic Bass model is also incorporated into the method of Meixell and Wu (2001), which consists of two phases. In the first phase, a PLC curve is constructed for a cluster of parts which are all considered to be similar. The PLC curve for the cluster is based on the usage pattern of a leading indicator, which is a product that is introduced slightly earlier than other products, but the PLC is similar, making the PLC of the product an indicator for other parts. Then, in the second phase, the deviations of the standard cluster PLC for a specific part are determined. This approach minimizes the number of possible scenarios that need to be investigated before determining the PLC curve, which should result in shorter computation times. Wu, Aytec, Berger and Armbruster (2006) describe in their article an approach similar to the first phase of the method.

One of the research areas within PLC forecasting focuses on the determination of transition points in the PLC. A transition point is the point in time a part goes from one stage in the PLC to the next stage. Within this area, articles are published by Kirshners and Sukov (2008).

A fourth stream of research focuses on forecasting a time interval in which a part is likely to become obsolete. Based on this information, strategies can be developed to manage the part as cost efficient as possible. This area is e.g. investigated by Solomon, Sandborn and Pecht (2000).

4.5 Which methods will be applied at SPO?

Based on the prerequisites set in section 4.1, the methods will be assessed to determine whether or not they are suitable for SPO. With respect to PLC clustering, the aim of the research is to cluster based on usage pattern. Historical usage will be used as a starting point, as is the case in the article of Meixell & Wu (2001), in which the clusters are created based on similarity of distribution parameters determined for a specific distribution and a specific set of historical usage. Although we will use historical usage as a starting point for clustering, we do not apply the method of clustering according

to distribution parameters. The main reason for this is, that we will start with clustering before determining which method to use, making it impossible to determine the parameter values first. Therefore we will use the observed patterns in the historical usage itself to determine the opportunities for clustering. Advantages of using historical usage are that this information is available over a time horizon of 5 years and it would give a clear indication of the usage patterns present, while being intuitively acceptable by the SPO team.

For PLC forecasting, historical usage data will be used, either from the product itself or from other products that appear to be similar. We do not set the requirement that the data should only be from previous generations, like Lapide (2002) described, but combine all parts with similar historical usage patterns.

Forecasts will be made based on historical usage according to both the Bass model (Bass, 2004) and statistical distributions, based on Petrescu (2009). The Bass model is a model in which forecasts are created based on three parameters, representing the probability of initial buy, the probability of a purchased due to imitator behavior and the scale of the sales. Different probability values lead to different curve shapes. Besides the Bass model we also selected statistical distributions to forecast the PLC shape. In his article, Petrescu mentions the Gamma distribution and he proposes a distribution called the Alpha distribution. Although the Alpha distribution might lead to accurate forecasts, we will not apply it. This is due to the requirement of simplicity and understandability, since standard distributions are known by the SPO team. A more extensive explanation of the selected methods will be provided in section 6.3.1 when they are applied in the forecasting model.

Both methods can be applied on both single products as on aggregated group usage, creating the opportunity to incorporate the created clusters. The combination of multiple approaches will give the opportunity to analyze the applicability of different forecasting models, which can be easily incorporated into simple applications. Furthermore, the different shapes the models can cover based on parameter changes fulfills the requirement that different types and lengths of usage patterns can be covered in the model.

4.6 Conclusion

For the model to be applicable at SPO, different requirements are determined. These requirements focus on the different shapes and lengths of the usage curves, product related data, applicability and understandability.

With respect to PLC clustering and PLC forecasting, different methods are described in literature. We selected a clustering method based on historical usage data. For PLC forecasting the basic principle of forecasting according to statistical distributions or the Bass model are considered to be the most promising approaches.

5 Which possibilities exist regarding product clustering using FRU characteristics?

Before being able to forecast the future usage of an FRU based on the usage pattern the FRU will follow, we need to know how the individual FRUs can be assigned to the usage patterns. Due to limited historical usage data, we cannot determine the usage patterns present, so therefore we will focus on assigning FRUs to the 12 partial usage patterns we have defined in section 3.4.3.

In the assignment of FRUs to partial usage patterns, the selection of the correct partial usage pattern should be based on aspects of the FRU that are known at the moment of the LTB, or characteristics, which can be determined for all FRUs in the product set. The desired result would be the possibility to make a well funded decision for a specific partial usage pattern or usage pattern based on the value of one or more characteristics.

To be able to determine the opportunities for product clustering, the first step is to determine the characteristics that might indicate a partial usage pattern. Once the characteristics are determined, we will test whether there is any relationship between the characteristics and the partial usage patterns based on statistical tests. The characteristics for which the statistical tests point in the direction of a relationship will be analyzed further with respect to applicability to identify partial usage patterns and the opportunities to cluster FRUs with similar usage patterns.

5.1 Which characteristics can be promising?

Characteristics are defined as aspects related to an FRU, known at the moment the LTB is executed and potentially influential to the remaining PLC of the FRU. In the selection of the characteristics that can be promising with respect to influencing the PLC of a FRU, a combination of characteristics from literature and expert opinion is applied.

5.1.1 Which characteristics are mentioned in literature?

In literature, different characteristics are mentioned and applied to forecast the PLC. Kirshners and Sukov (2008) use similar product data regarding the time the product was on the market. The time on the market could be an influential characteristic, since the longer the time the FRU is on the market, the longer the PLC will be. In a long PLC, possible shapes are a long tail or fluctuation in usage. These effects might be less visible in short PLCs. The age of the FRU is an indicator for spare part usage according to Lapide (2002). Here we will define age as the number of years between the introduction and LTB date, and can be used as an indication of the PLC stage the FRU could be in, in combination with the time on the market.

Meenaghan and O'Sullivan (1986) mention aspects that can influence the shape and length of the PLC. Two of the aspects mentioned in their article are technical problems and innovation. These issues can cause an abrupt start or stop of the PLC for a single FRU, with usage continuing in the pattern of another. When the FRUs are combined in one PLC, the effect on the final PLC will most likely be minimal. However, when assessing the PLC pattern for the single FRU, the resulting PLC will be different. Meenaghan and O'Sullivan also mention price, brand, recycling possibilities and technology. The price can influence the amount of purchases, as could the brand in which the FRUs are used. Recycling possibilities might lengthen the PLC, because supply is continued for a longer period, while (changes in) technology will determine the time period in which the FRU can be used.

Forecasted reliability is used as an estimator of replenishment forecasting by Singh and Sandborn (2006). Although this article focuses on the development of replacement schedules for maintenance, forecasted reliability might be an indicator for the PLC of the FRU, since low reliability is expected to lead to high usage. Finally, Rink and Swan (1979) mention the level of product aggregation as characterization criterion. Different levels of aggregation can result in different PLCs that can be used as a standard, assuming that all FRUs have a similar PLC. Within SPO, the levels of aggregation that can be distinguished are aggregation on brand, commodity and division.

5.1.2 Which characteristics are mentioned by experts from the SPO team?

Within the PLCM team, the RSP is mentioned as a characteristic. The idea exists that the longer the RSP, the more difficult it is to forecast the required LTB need. This is based on the fact that the longer the time horizon that needs to be forecasted, the more the usage pattern can deviate from what is expected.

Besides the RSP, the TM144 status is added. The TM144 status can either be 'pre' or 'post', indicating whether the LTB takes place respectively before or after production stop. If the TM144 status is pre, there is a high probability the FRU has not yet reached the peak in spare part usage, while the peak is considered to be passed when the status is post.

Finally, the month forecast used in the LTB calculation is added as possible influential characteristic. If the LTB Month forecast is high, the PLC is expected to have a different shape then when the LTB month forecast is low. For example, when the LTB month forecast is less than 1, the expected usage pattern is more likely to be low and intermittent than a standard PLC curve.

5.1.3 Which characteristics are considered to be interesting?

We derived a set of 13 characteristics based on literature and expert opinions. Although these 13 characteristics are valuable in the context of the research they were used in, this does not implicate that the characteristics are also interesting to investigate as possible indicator of a specific usage pattern at SPO. We will determine the applicability in this section, together with a clear statement of how to measure the characteristics, to eliminate inconsistencies due to incorrect measurement.

From the 13 characteristics under consideration, 7 are considered to be irrelevant for the situation of SPO. First, the characteristic Time on the market is excluded, because the value of this characteristic is equal to the sum of the Age FRU and the RSP. The decision is made to select the last 2 characteristics, because they can be used to get an indication of the point in time the LTB is executed, which might be an indication for the PLC stage the FRU is in. The characteristics Technical problems, Innovation and Technology are excluded based on the absence of data that can be used to accurately determine the characteristic values. Price and recycling possibilities are excluded because they do not have a direct influence on the PLC. For customers of SPO, price is not an important factor in acquiring a spare part, because most customers have a service contract and pay a fixed fee per period. Recycling possibilities is considered to be a supply problem, which will not affect the shape of the PLC.

A consistent way of measuring the characteristic values provides the opportunity to reproduce the values, and will result in consistent and comparable calculations for different FRUs. To be able to measure the characteristic values in a consistent manner, we determined a way of measurement for each of the characteristics under consideration. The way of measurement can be found in Table 5-1.

Characteristic	Way of measurement
Age FRU	Years between introduction and the LTB date (rounded to integers)
Brand	The brand to which the FRU belongs
Commodity	The commodity the FRU belongs to
Division	The division the FRU belongs to
Forecasted reliability	The theoretical removal rate, determined at the introduction of the FRU
LTB Month forecast	The forecasted monthly amount at the moment of the LTB
Remaining Service Period	Years between the LTB date and the EOS date (rounded to integers)
TM144 Status	Pre when the LTB is executed before production stop, else Post

Table 5-1: Characteristics and their way of measurement

5.1.4 Conclusion

We have derived a list of 13 potential characteristics that might influence the PLC from literature and expert opinions. Based on applicability to the case of SPO and data availability, the list is reduced to a set of 8 characteristics that will be used in the remainder of this chapter. These 8 characteristics and the way in which they are measured are described in Table 5-1.

5.2 Which characteristics influence the PLC of the FRU?

In our research we aim to identify usage patterns and forecast future usage based on the remaining usage pattern. Since we do not have the full usage patterns, here we will focus on the existence of a relationship between the characteristics and the partial usage patterns. If this relationship exists, than changes in the characteristic values can be used as an indication for a change in the partial usage pattern, and combining these partial usage patterns will provide insight in the overall usage pattern.

Before being able to assign an FRU to a partial usage pattern based on characteristic values, we first need to establish whether there is actually a relationship between the characteristic and a partial usage pattern. In case there is a relationship, the (combination of) characteristic values for a specific partial usage pattern should be similar for all FRUs with that pattern, but substantially different from the characteristic values of other patterns. This creates the ability to assign an FRU with that specific (combination of) characteristic values to the partial usage pattern. For example, if the age of the FRU for FRUs belonging to the increasing usage pattern is always 2 years, while it is more than 5 years for all other partial usage patterns, we can assign an FRU with the age of 2 to the increasing pattern without hesitation.

To determine if there is a relation between the characteristics determined in section 5.1 and one or more of the partial usage patterns, a two step approach will be applied. In this approach, we will use the assumption that the characteristic values will substantially differ between the partial usage patterns, indicating they are related to a specific partial usage pattern. We will perform statistical tests to see whether this is actually the case and for which characteristics. As a first step, in section 5.2.1, statistical tests will be executed over the model building selection to determine whether there is a relation between the characteristic value range and the partial usage patterns. Section 5.2.2 represents the second step, in which the tests will be repeated for a selection of the 8 largest commodities with respect to the LTB value, to determine whether the results of the overall dataset comply with the results for a specific subsection of FRUs. For both sections, detailed results can be found in appendix F.

5.2.1 Is there a relationship between the characteristics and partial usage patterns?

In the first step of the statistical test to determine relationships between characteristics and partial usage patterns, we will focus on the relationship between the entire value range of the characteristic and all 12 partial usage patterns. The outcome of this test can be used to determine whether the characteristics can actually be used as an indicator of a partial usage pattern, because if there is no relation between the characteristic and the partial usage patterns for a large set of FRUs, the probability of a relationship between a specific characteristic value and one of the partial usage patterns is unlikely. As a result, all characteristics for which we cannot find a relationship in this section will be omitted from the research.

Before we are able to perform a statistical test, we should first determine which test to apply. The selection is made based on the types of variables, because different statistical tests are designed for different variable types. In general there are 4 different variable types, being:

- Nominal (qualitative, multiple options, cannot be ordered, like color)
- Ordinal (qualitative, multiple options, can be ordered, like answer to the question 'How happy are you?')
- Interval (quantitative, like temperature)
- Ratio (quantitative, zero means no value, like distance)

Since we want to test the relationship between the partial usage patterns and the characteristics, the independent variable will be the partial usage patterns, which is considered a nominal value because the patterns cannot be ordered. As a result, we can only select statistical tests which are developed for at least one nominal variable, but the selection of the specific test depends on the dependent variable type, determined by the 8 selected characteristics. The characteristics Age FRU, RSP and Forecasted reliability are interval variables, because an age of zero means the introduction date is less than a year ago, not that there is no age yet. The same holds for the RSP. The LTB Month forecast is a ratio variable, because here a forecast of zero means there is no forecast. Both interval and ratio variables, which are not normally distributed, combined with a nominal variable, can be tested using the Kruskal-Wallis approach (Dodge, 2008, pp. 288-290). This test can be used for testing the hypothesis that rankings are the same in different groups. In these cases, the hypothesis that will be tested is that the average value of the characteristics is equal for all partial usage patterns.

The characteristics Brand, Commodity, Division and TM144 Status are all qualitative variables that cannot be ranked. Therefore they are categorized as nominal variables. A combination of two nominal variables can best be tested using the Chi-Squared test of Independence (Dodge, 2008, pp. 79-81). This test can be used to determine the probability that the observed values are found without any connection between the variables. If there is a small probability the results could have been found without a connection, the characteristic might be influential.

To perform the tests, we used the program SPSS Statistics, which contains both the Chi-Square test of Independence and the Kruskal-Wallis test. In both tests, the hypothesis will be rejected when the results are considered statistically significant, which is the case when the p-value is smaller than 0,05. This indicates that the probability that the observed values could have occurred at random, so without any connection between the two variables, is less than 5%.

After conducting the tests, the results indicate that for all 8 characteristics the p-value is smaller than 0,001. This indicates that the probability the results are created without any connection between the variables is very small. Based on this result we can state that there is a relationship between the characteristics and the partial usage patterns on an aggregated level.

5.2.2 Is there a relationship between the characteristics and partial usage patterns on commodity level?

Now we have established that there is a relationship between the 8 characteristics and partial usage patterns on a general level, the second step is to determine whether all 8 characteristics are related to the 12 partial usage patterns for a specific subset. Here the subsets under consideration are the commodities, because within SPO the assumption is that FRUs in a commodity follow the same usage pattern. If characteristics are only related to a specific commodity, this might also indicate that the characteristic only relates to a specific usage pattern. This knowledge increases the opportunity to correctly assign an FRU based on a characteristic value, because the amount of characteristic value combinations to assess decreases.

The statistical tests will be executed for the commodities with a total LTB value exceeding €1 million. Because the variables do not change, we do not need to reconsider the variable types. However, we do need to reconsider the statistical tests to apply, because the amount of cases under consideration have decreased, which could lead to the inapplicability of statistical tests. For the Kruskal-Wallis test, samples sizes do not matter, so we can apply this test again. The Chi-Squared test of Independence requires a minimal sample size of 5 observations per combination of values, which will be violated due to the limited amount of FRUs belonging to the commodities. As a result, we will use Fisher's Exact test (Dodge, 2008, pp. 205-206), which will also test the probability of getting the observed results without any connection between the variables, and is especially suitable for small sample sizes. Due to the limits in computing Fisher's Exact test using SPSS Statistics software, we will compute the values for combinations of 2 partial usage patterns and 2 characteristic values, while making sure all possible combinations are investigated. If for more than 50% of the characteristic values at least one combination has a p-value smaller than 0,05 we consider a relationship between the characteristic and the partial usage pattern as likely.

After conducting Fisher's Exact test, the results for the specific commodities differ from the general results. One of the differences is, that there is no relationship between the characteristic Division and the partial usage patterns within the commodity TAPE. For all other commodities, this relationship exists, but in most cases only for a limited set of partial usage patterns. This could indicate that the characteristics could only be related to a selection of the partial usage patterns. The same conclusion can be drawn for the characteristic Brand in the commodities POWR and TAPE. For the characteristic TM44 Status, the commodity PLNR appears to be the only commodity for which this characteristic might have a relationship with the partial usage patterns, but the commodity MISC might also be an interesting option to consider.

When we analyze the results of the Kruskal-Wallis results for the selected commodities, we also see different results than was the case for the general selection. For all commodities, the LTB Month forecast still shows significant results, indicating that this could be an influential characteristic for all commodities. The characteristic RSP only gives significant results for the commodities CARD, CPU, MEM, PLNR and POWR. When we look at the characteristic Age FRU, the results also indicate a

limited set of commodities for which this characteristic could be influential, namely CARD, HDD, PLNR, POWR and TAPE. Finally, the Forecasted reliability shows promising results for the commodities CARD, CPU and MISC. Concluding we can state that potential characteristics differ per commodity, and even per subsection within a commodity.

5.2.3 Conclusion

Based on the statistical test results, we can conclude that all 8 characteristics under consideration are likely to be related to at least one partial usage pattern. Furthermore, the characteristics that could be related to a specific partial usage pattern differ per commodity and sometimes even per partial usage pattern within a commodity. It is also not yet possible to determine which characteristic is most likely to provide good results with respect to identifying the partial usage pattern, since the results differ based on the grouping approach under consideration. As a result, the characteristics Age FRU, Brand, Commodity, Division, Forecasted reliability, LTB Month forecast, RSP and TM144 Status are all considered to be potential characteristics for identifying partial usage patterns.

5.3 Which characteristics can identify the partial usage patterns?

Based on the statistical tests, we concluded that all 8 characteristics could potentially identify a partial usage pattern for all commodities or just for a specific subsection. However, this knowledge does not provide the opportunity to actually identify a partial usage pattern yet. To be able to identify a partial usage pattern, a (set of) characteristic value range(s) needs to be established that correctly identifies a specific partial usage pattern. This would then create the opportunity to assign all FRUs that have the specific set of characteristic value ranges to a partial usage pattern. Another opportunity is to create guidelines according to which a limited set of possible partial usage patterns can be used as a starting point for the assignment process.

Because we focus on characteristic value ranges, the first step will be to determine whether the value ranges for the characteristics are correctly representing the dispersion of the values belonging to the characteristics, or if outliers create an incorrect image of the value range. We identified 2 outliers, as described in appendix G. After removing the outliers, we will focus on the possibility to identify a specific value range that corresponds to one of the partial usage patterns.

Due to the high amount of FRUs under investigation, we decided to focus on the commodity HDD in this section, to be able to do a more thorough analysis of a limited set of FRUs that is considered to be similar. Since we focus on the commodity HDD, a selection of the 8 identified characteristics is excluded. The characteristics Forecasted reliability, RSP and TM144 Status are already classified as not influential based on the statistical tests, and will therefore not be incorporated. And because we focus on a single commodity, we also exclude the characteristic Commodity, since this characteristic can only take on one value in our selection. As a result, we will focus on the characteristics Age FRU, Brand, Division and LTB Month forecast.

First we will focus on the characteristic Division. Due to differences in grouping, the LTB division eServer consist out of a combination of the IBM divisions Mainframe and Power, as where described in appendix A. Furthermore, two LTB divisions have a different name. The part of the IBM division Power that is not incorporated in eServer is called ITS, and the IBM division System X is referred to as xSeries. Because it is difficult to reallocate the FRUs from the combined LTB division to the IBM divisions, we will use the LTB divisions eServer, ITS and xSeries.

When we analyze the results of Fisher's Exact test for HDDs, a number of significant differences can be identified. The first is, that the division eServer has a relatively large amount of FRUs that show low usage patterns compared to the other divisions. Another difference is that the Storage division has a higher percentage of FRUs with a fluctuating pattern. The no usage pattern appears to be most common for the FRUs in the division xSeries. If we test these conclusions on the HDD test selection, consisting of 47 FRUs, the first observation of a relatively large amount of low usage patterns for the division eServer could be confirmed. The other two conclusions cannot be confirmed based on the test selection, since the amount of FRUs with a fluctuating pattern is almost similar to the amount for other divisions and there are no FRUs with no usage in the test selection for the division xSeries.

The next characteristic under consideration is the characteristic Brand. In total there are 8 different brands, of which 7 contain HDDs. When we analyze the results of Fisher's Exact test for Brands, the brand Disk products appears to have approximately 75% of the FRUs almost equally divided over 6-20 demands, slow decline, fluctuating and strong decline. For System p, most FRUs have either relatively low usage or show a slow or strong decline pattern. When we test these assumptions on the HDD test selection, the results do not confirm nor disprove the assumption.

Since the results of Fisher's Exact test for both Division and Brand are not confirmed in the analysis of the test selection, it feeds the thought that there is no distinct relationship between the partial usage patterns observed and the Brand and/or Division the FRUs belong to. All partial usage patterns could be observed in all situations. Although the number of times the specific patterns are observed might differ, the differences are not clearly pointing in the direction of a (set of) specific partial usage pattern(s) that could be assigned to a Brand or Division.

The other characteristics under consideration, being Age FRU and LTB month forecast, are analyzed using the Kruskal-Wallis method in the previous section. As a result, only overall results are available, and additional actions need to be executed before being able to determine which differences can be identified.

To be able to distinguish whether there are differences between the partial usage patterns, we will create box plots to see the range of characteristic values belonging to a certain partial usage pattern. Preferably, all value ranges are non-overlapping, which creates the opportunity to assign a characteristic value range to a specific partial usage pattern. An example of the preferred box plot result is represented in Figure 5-1.

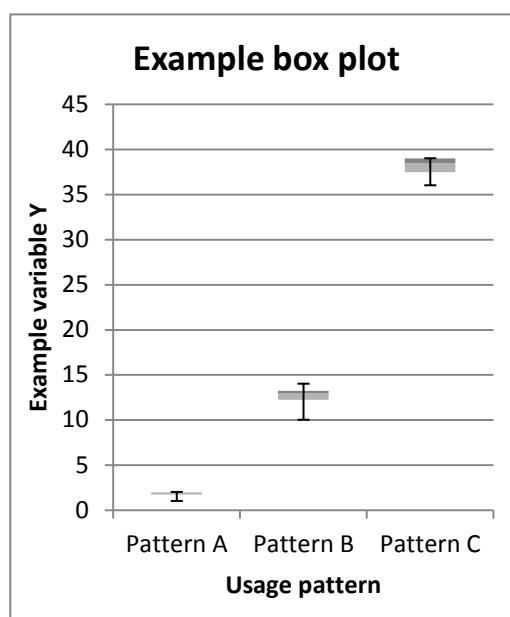


Figure 5-1: Example preferred results Box plots

In Figure 5-2 the box plot for the LTB Month forecast is displayed. Especially for the LTB month forecast, there is a large range between the minimum and maximum forecast value, with a substantially smaller range for the first three percentiles. These large ranges create a difficulty in determining which forecasted values could be related to specific partial usage patterns, due to substantial overlap.

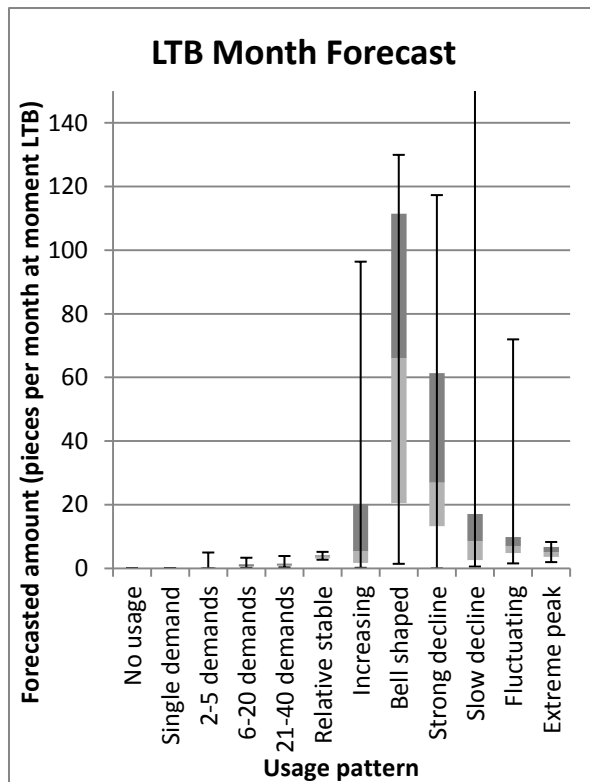


Figure 5-2: Box plot LTB Month forecast

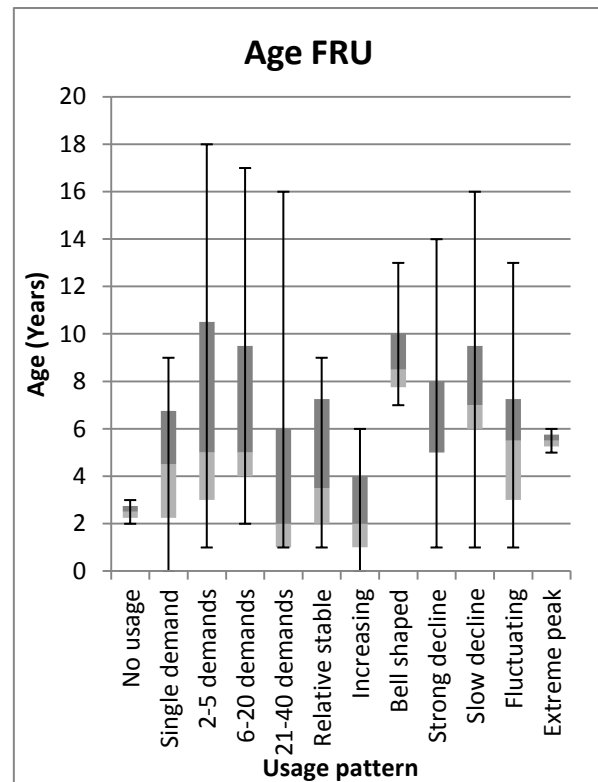


Figure 5-3: Box plot Age FRU

A clear difference can be identified between the more irregular patterns with an LTB Month forecast of less than 5 pieces, represented by the patterns named '... demands' and relative stable, and the remaining partial usage patterns with a more constant usage and a higher forecast.

When we compare the results of the model selection with the results of the test selection, which can be found in appendix H, the large ranges for the bell shaped, slow decline and fluctuating patterns are similar, as are the low forecast values for the more irregular patterns. The forecasted value for the extreme peak, which reaches to a value of more than 100, differs largely from the result in the model building selection. For the increasing pattern, the forecasted value is much lower than in the model selection, while for the strong decline the range is smaller. Based on these results, we can conclude that the LTB Month forecast value might be used as an indicator to narrow down the possible partial usage patterns, but cannot be seen as an indicator that accurately identifies the actual partial usage pattern a specific FRU follows.

When we analyze the box plot for Age FRU, which can be found in Figure 5-3, the differences are less clear between the different partial usage patterns. This is not as one would expect when the traditional PLC pattern is applicable to this situation. In that case, the increasing pattern would show a lower age range than the decreasing patterns, but from Figure 5-3 we can conclude this is not necessarily the case, because of the overlap in ages for both partial usage patterns. This could be a misleading view if the expected times on the market differ largely, which could lead to an overlap in the increasing phase for an FRU with a long time on the market and a decreasing pattern for FRUs with a short time on the market at a certain point in time. But since we have only incorporated HDDs into the selection, we would expect a similar time on the market, without overlapping phases. The similarities in age ranges for all partial usage patterns are also confirmed in the test selection, as can

be seen in appendix H. As a result, we conclude that Age FRU is not a straightforward indicator for partial usage patterns.

Concluding we can state that, based on the statistical tests and the comparison of the results between the model building selection and the test selection, there is no clear indication that any of the potentially influential characteristics will give a representative indication of the partial usage pattern an FRU might follow on a standalone basis. It is only possible that single characteristics give an indication of the potential partial usage patterns the FRU might follow. Furthermore, it appears that characteristic value ranges cannot be assigned to a specific partial usage pattern, because similar characteristic values can be found for different partial usage patterns.

5.4 Can a combination of characteristics identify partial usage patterns?

We analyzed the influence of the characteristics, and concluded that none of the characteristics in the analysis is able to give a clear indication of the partial usage pattern individually. In this section, we investigate whether a combination of two or more characteristics is able to identify a specific partial usage pattern. The combinations are created based on the full characteristic value range, partial characteristic value ranges and historical usage. To be able to compare the results, we again focus on the HDD selection and the characteristics Age FRU, Brand, Division and LTB Month forecast. We will focus on forecasting FRUs with an increasing usage pattern, since this partial usage pattern represent a small subsection of the HDD selection, which makes them more difficult to identify.

5.4.1 Identification based on combinations of full characteristic value ranges

To determine whether a combination of full characteristics value ranges can be related to a specific partial usage pattern, we determine the value range for all characteristics corresponding to the increasing usage pattern. In the identification process, different combinations of the characteristic value ranges will be applied, and the percentage of successfully identified HDDs with an increasing pattern will be determined. If a certain combination results in a selection that consists for 80 percent or more out of the desired partial usage pattern, the combination is characterized as promising and will be used on the test selection, to see whether similar results are realized on a different group of FRUs.

Before being able to start, information about the data set under consideration and the value ranges of the characteristics is required. With respect to the dataset, the complete HDD selection consist of 152 FRUs, of which 11 FRUs are categorized as FRUs with an increasing usage pattern. The value ranges of the characteristics for the increasing pattern are represented in Table 5-2. For these values we executed an outlier analysis, but no outliers are found, as described in appendix G.

Characteristic	Value range increasing pattern
Division	Storage, eServer, ITS, xSeries/PWS, RSS
Brand	Disk products, System p, Power Systems, Server, Point of Sale
Age FRU	<1 – 6
LTB Month forecast	0,19 – 96,42

Table 5-2: Characteristic values increasing pattern HDD

When we analyze the results in Table 5-2, it becomes clear that for the characteristic Division all possible characteristic values are represented in the selection of FRUs with an increasing usage pattern. As a result, using this characteristic in the clustering process does not add any value, since the selection will not change. For the characteristic Brand, 5 out of the 7 possible brands are

represented in the range. This does not indicate that for the remaining 2 brands increasing usage patterns do not exist, but they are not included into the selection. The characteristic Age FRU covers a selection of FRUs that are relatively young, since expected times on the market can be up to 18 years. The range for the LTB Month forecast is considered to be very large.

Now the characteristics under consideration and their corresponding value ranges are known, all possible combinations of the characteristics will be used to determine the amount of FRUs compliant to the value range. Since we only use value ranges compliant to the value range of the 11 FRUs with an increasing usage pattern, the closer we will reach a total amount of 11 FRUs selected with the range, the better the combination is in identifying FRUs with increasing usage patterns.

Based on the analysis of the results for the entire range, the best performance is realized by a combination of the characteristics Brand, LTB Month forecast and Age FRU. This leads to a total amount of 70 FRUs which comply to the value ranges, of which 11 show an increasing usage pattern. Although this is the best result, it is not a result that gives a clear indication of FRUs with an increasing usage pattern, since only 16% of the selection has this pattern.

To determine whether similar results are realized for other partial usage patterns, we repeat the process for a combination of the slow and strong declining usage pattern. We selected these 2 patterns because the trend is completely opposite to the trend in the increasing usage pattern. When it is also not possible to assign a specific value range to the declining usage patterns, there is a high probability this is also not possible for any other partial usage pattern.

There are 52 FRUs with one of the two declining usage patterns. The best result for the combination of characteristics is 136 FRUs for the same combination as for the increasing usage pattern, resembling 38% successfully identified. Although the correctly selected percentage is higher for the declining usage patterns, it is still less than 50%.

Since the results for both partial usage patterns do not reach the boundary value of 80% correctly identified, this indicates that the combination of the complete value ranges does not lead to a high probability of identifying FRUs with a specific partial usage pattern, and we can conclude that the full range covers too many values to be able to accurately identify a partial usage pattern.

5.4.2 Identification based on combinations of partial characteristic value ranges

As we concluded in the previous section, the full characteristic value range is too wide, leading to inaccurate results when trying to assess the FRUs that follow a specific partial usage pattern. In this section, we will divide the characteristic value range over a number of subsections. To be able to compare the results, we will focus on the increasing usage pattern, using the same characteristics as in the previous section.

To determine the sub ranges for the characteristics Age FRU and LTB Month forecast, we will use the percentile values as potential boundary value. Using different combinations of the minimum, 25 percentile, median, 75 percentile and maximum value, 10 different sub ranges can be defined. The final sub ranges are determined such that the probability of identifying FRUs with the increasing usage pattern is maximized. The resulting sub ranges are presented in Table 5-3. For the characteristic Brand, the sub ranges are determined by the different brands under consideration.

Characteristic	Sub range 1	Sub range 2	Sub range 3	Sub range 4	Sub range 5
Age FRU	<1 – 1	1 – 2	2 – 4	4 – 6	
Brand	Disk products	System p	Power systems	Server	Point of Sale
LTB forecast	0,19 – 1,72	1,72 – 5,4	5,4 – 96,42		

Table 5-3: Selected sub ranges increasing usage pattern

We will perform a factorial analysis to determine the amount of successfully identified FRUs using combinations of the characteristic sub ranges. Due to the large amount of possible combinations, we only select the options for which the percentage successfully identified FRUs is larger than the boundary value of 80 percent. In total, 6 combinations of sub ranges exceeded the boundary value and will be applied on the test selection. The combinations of sub ranges, the amount of FRUs selected with this sub range and the percentage of the selection correctly identified can be found in Table 5-4. In the most right column of Table 5-4, the results when applying the sub ranges to the test selection are displayed. The HDD test selection consists out of 46 FRUs, of which 3 have an increasing usage pattern. To clarify the way the results in the most right column should be interpreted, we will use option 1 as an example. For option 1, only 3 of the 46 FRUs from the test selection are selected based on the defined characteristic value ranges, of which 67% (or 2 FRUs) have an increasing usage pattern.

Option	Age FRU	Brand	LTB Month forecast	Results model selection	Result test selection
1	<1 – 1	Disk products	-	1 selected, 100% correct	3 selected, 67% correct
2	1 – 2	Disk products	5,4 – 96,42	1 selected, 100% correct	1 selected, 0% correct
3	1 – 2	Power systems	-	1 selected, 100% correct	1 selected, 100% correct
4	2 – 4	Point of Sale	-	1 selected, 100% correct	1 selected, 0% correct
5	2 – 4	Power systems	-	2 selected, 100% correct	0 selected, 0% correct
6	4 – 6	Power systems	-	1 selected, 100% correct	1 selected, 0% correct

Table 5-4: Results identification of increasing usage pattern with characteristic sub range combinations

The smaller sub ranges lead to the identification of three FRUs with an increasing usage pattern in the test selection, using sub range combinations 1 and 3. However, the selected amount of FRUs is very small in all cases, making it difficult to assess whether this results will also be realized in a large data set. Therefore we execute the same procedure for the selection of FRUs that have a slow or strong declining usage pattern. Of the sub ranges under consideration, 23 options result in a selection that consists for 80% or more out of FRUs with a slow or strong declining pattern. When we test these options on the HDD test selection, only 7 reach the boundary value of 80% successfully identified. And as with the increasing usage pattern, the amount of FRUs identified is limited.

The small amount of FRUs identified with each sub range combination for the partial usage patterns has a large disadvantage. In this case, we only determined sub range combinations to identify FRUs with an increasing or declining usage pattern, and we need more than one sub range to identify the FRUs with those partial usage patterns. If these two partial usage patterns were the only patterns to identify, this would not be an issue. However, there are 12 different partial usage patterns. If all partial usage patterns would require multiple options to identify the FRUs with those patterns, we will end up with a large set of characteristic sub range combinations. If the PLCM staff has to check all these potential options, the workload will increase, which is not desirable.

Based on this analysis, we conclude that using sub ranges for the characteristic values results in a more accurate identification of partial usage patterns, but the percentage of successfully identified FRUs remains low. A drawback in this approach is the large amount of sub ranges that are required to identify each partial usage pattern. Therefore we conclude that, although the results of the combinations of sub ranges are more accurate, the results do not compensate the large amount of computational effort required. As a result, identifying partial usage patterns based on combinations of sub ranges is not considered to be an appropriate identification method.

5.4.3 Identification based on historical usage data

From the previous sections, we concluded that the characteristic values cannot be used to actually identify the partial usage patterns efficiently. An additional indicator in the identification of partial usage patterns might be historical usage. If the historical usage shows a declining trend, a partial usage pattern with a declining section is more likely than an increasing usage pattern if both options are possible.

To limit the amount of additional data, 4 historical usage points are added. The LTB date is selected as reference point, and the usage added is the usage in 1, 3, 5 and 7 periods before the LTB date, in which a period is defined as a 4 week time bucket. The selected periods have one period in between, to be able to cover a larger time horizon, which might make the underlying trend easier to identify, while still limiting the amount of additional data.

One of the problems we encounter while adding the historical usage to the HDD modeling selection is, that usage data is not always available for 7 periods back, due to limited data availability. To be able to determine whether adding historical usage is valuable, we exclude the FRUs with partial usage data available. The remaining selection consists of 99 of the 162 FRUs originally selected.

When we analyze the dataset, we first focus on the 17 FRUs that show zero usage in the selected usage periods. These FRUs all belong to one of the partial usage patterns no usage, single demand, 2-5 demands, 6-20 demands or 21-40 demands. For the 5 FRUs which have usage in only one of the previous periods, the results also point into the direction of one of these partial usage patterns. When the amount of periods in which usage has occurred exceeds 1, the intermittent usage pattern becomes exceptional. Based on this observation we state that having no or only a single period with historical usage most likely points toward an intermittent usage pattern.

Part	Usage 7 period ago	Usage 5 periods ago	Usage 3 periods ago	Usage 1 period ago
1	0	1	0	2
2	0	0	1	1
3	0	1	1	3
4	14	21	12	16
5	95	81	65	105
6	22	13	23	18

Table 5-5: Historical usage parts with increasing pattern

For the FRUs which have historical usage in 2 or more periods, the results are less straightforward. Table 5-5 shows the historical usage of FRUs that are categorized as FRUs with an increasing usage pattern based on the usage in the entire time window. For an increasing usage pattern, we would expect the historical usage to increase when going from left to right in the table. For the first three FRUs in Table 5-5, a slow increase could be an option, but constant or intermittent usage could also

be seen in the pattern. For part 4 and 6, the idea of a fluctuating pattern could easily be developed based on the constant increase and decrease in usage, without any clear trend. Finally, the usage of part 5 is decreasing over the first periods, which could lead to discarding the idea of increasing usage.

For other partial usage patterns, the historical usage shows similar results, with a trend in historical usage that does not correspond to the partial usage pattern one would expect based on the complete horizon. Since we cannot know future usage at the moment an LTB is executed, using limited historical usage data to determine the partial usage pattern might lead to less accurate forecasts for the usage over the RSP. Using a complete historical usage pattern might provide better results, but this cannot be tested due to limited data.

5.4.4 Conclusion

In the previous subsections different approaches are investigated, all focused on identifying the partial usage pattern an FRU might follow, based on which FRUs with similar partial usage patterns can be clustered.

The first approach focused on identification using a combination of full characteristic value ranges. The results for both the increasing and declining usage pattern indicate that combinations of the full characteristic value ranges do not lead to a high probability of identifying FRUs with that specific usage pattern. We concluded that the same characteristic values are found for different partial usage patterns, and therefore full characteristic value ranges cannot be used as a reliable indicator for a specific partial usage pattern.

In the second step, smaller characteristic sub ranges are selected. Based on the results of a factorial analysis, a small selection of sub range combinations was determined, for which positive results were realized for a small set of FRUs. Although this is a promising result, a large set of sub ranges is required to cover the entire selection of FRUs, which leads to the conclusion that the computational effort is larger than the benefits that could be realized. The investigation of the added value of historical usage in identifying the partial usage pattern provided contradictory indications, since the trend in the historical usage data did not correspond to the observed usage pattern over the entire time window.

Based on these results, we conclude that characteristic values are not related to a specific partial usage pattern, but to the FRU itself. As a result, they cannot be used to accurately identify a specific partial usage pattern, but can only be used to give an indication of the partial usage patterns the FRU might follow, based on predetermined value ranges.

5.5 Can characteristics be used to identify fast and slow moving FRUs?

Based on the LTB forecast ranges and the identified partial usage patterns, it is likely that both fast and slow movers are present in the selection of FRUs from SPO. According to Silver et al. (1998, p. 318) slow movers are defined as parts that have a demand of less than 10 pieces in the lead time of the product. Furthermore, Silver et al. (1998) mention different approaches specifically for slow moving parts. We will try to identify slow movers, using the characteristic information gathered from the previous sections. Because we do not know the lead time, we will use the assumption that FRUs are considered slow movers when the usage is less than 10 pieces a year. Because a lead time of a year is rare, this approach provides the certainty that when an FRU is identified as slow mover based on yearly usage, it will also be a slow mover based on lead time usage. When FRUs that could be

considered as slow movers can be identified, it will also be possible to cluster slow movers in a single group. For this group we can then assess whether the forecasts made with the current forecasting methods are relatively accurate, or if specialized forecasting methods for slow movers might lead to more accurate forecasts.

In the identification process, we will use the assumption that all FRUs with one of the following partial usage patterns is considered a slow mover: no usage, single demand, 2-5 demands, 6-20 demands or 21-40 demands. The reason these partial usage patterns are selected is that there is a large probability that the annual usage is less than 10 pieces a year over the 5 year timeframe we consider. In the identification process, we will focus on the LTB Month forecast and historical usage. The LTB Month forecast is selected because a low month forecast is considered to be related to low usage, and historical usage is selected based on the results of section 5.4.3, in which we state that having no or one period with historical usage is often related to an intermittent usage pattern.

The first step is to identify the LTB month forecast range which covers most of the FRUs with one of the selected partial usage patterns, based on the box plot values determined in section 5.3. Since the maximum value for 4 of the 5 partial usage patterns is below 4, and the third quartile for the fifth partial usage pattern is also below 4, we will select a forecast range from 0 to 4 as a starting point. When we test the LTB Month forecast range on the model building selection, 20 out of the 22 FRUs with one of the desired partial usage patterns are included, representing a coverage of over 90% for the model building selection. When we apply the LTB Month forecast range on the test selection, all 8 FRUs with the selected partial usage patterns are correctly identified. Based on this result, we conclude that the defined range for the LTB Month forecast can be applied to identify the FRUs with the preferred partial usage patterns.

The second step focuses on selection based on historical usage in none or 1 period. As in section 5.4.3, we will use the 99 FRUs for which usage data is available for all selected periods in the model building selection. After applying the LTB Month forecast range, there are 54 FRUs left, including the 22 FRUs which have the selected partial usage patterns. As a second step, we will select the FRUs which have usage in 0 or 1 of the 4 historical usage periods selected. This results in a set of 22 FRUs, of which 21 have one of the desired partial usage patterns, representing 95% successfully identified. As a reference, we apply the same approach on the test selection. The test selection consists of 25 FRUs for which all historical data is available, with 15 FRUs having an LTB Month forecast in the range from 0 to 4, of which 8 have one of the desired partial usage patterns. After the selection based on the number of periods with historical usage, 8 FRUs remain, including 6 with one of the selected partial usage patterns. This represents a total of 75% of the FRUs successfully identified. Although the percentage is lower for the test selection, it is close to the 80% we specified as boundary value.

Based on the results, we conclude that the combination of a LTB Month forecast value range from 0 to 4 and a maximum of 1 period usage for the historical usage periods selected is an appropriate method to identify slow moving FRUs.

5.6 Can the characteristics identify the usage pattern and PLC stage?

Based on the results of sections 5.2 to 5.4, we concluded that the selected characteristics based on literature and expert opinions cannot be used to identify the partial usage pattern of an FRU, they can only give an indication of possible partial usage patterns. This is not the desired and expected outcome, because being able to identify partial usage patterns would simplify identification of the

usage pattern as a result of the changing characteristic values over time. However, identifying the usage pattern based on partial usage patterns might not be the only possible approach. In this section, we will try to construct a full usage pattern and determine whether characteristic values can be related to this usage pattern. When possible, this approach would create the opportunity to forecast the LTB need based on the usage pattern more accurately while reducing excess stock.

One of the difficulties in this approach is that we do not have usage data for a full usage pattern. Therefore we created a usage pattern for the commodity HDD by reallocating the historical usage to make sure the PLC stages are aligned and calculate the correlation between the observed usage. If the calculated correlation between the reallocated usage for different FRUs is high, the overlapping section of the usage is considered to follow a similar pattern. The characteristic values will then be analyzed, to see to which extent they are similar. The basis assumption underlying this approach is that FRUs in a commodity are expected to have the same usage pattern.

The correlation between the overlapping periods of usage is determined according to the following formula:

$$\rho(X, Y) = \frac{cov(X, Y)}{\sigma(X) * \sigma(Y)} \quad (8)$$

with X being the usage for the first FRU in the calculation and Y the usage for the second FRU in the calculation. The amount of periods over which the correlation is calculated, is determined based on data availability. If the reallocated usage for part X covers the periods 20 to 60 and the reallocated usage for part Y covers the periods 32 to 72, then the usage to determine the correlation on is the usage in the periods 32 to 60, since for those periods usage is available for both FRUs. A minimum of 5 overlapping periods is taken into account, to prevent high correlations based on a few periods.

When we assess the correlation between the usage of the 2 FRUs, the results can differ between -1 (indicating the usage patterns are contradictory) and 1 (indicating the usage patterns are similar). To construct the complete usage pattern, we will start with an FRU for which information from the beginning of the PLC is available. We then determine the correlation of this FRU with others, and when the correlation coefficient of the usage of 2 FRUs is 0,7 or higher, both are included in the usage pattern. This process is executed for all FRUs that are already included. It is not necessary that all FRUs have high correlations between them, because correlations cannot be calculated for all combinations due to differences in available usage periods. If mutual correlation for all FRUs would be considered as a requirement, the usage pattern could only be created based on complete PLC data, which is not available, making it an impossible requirement to take into account.

Based on the described procedure, 30 of the 152 HDDs are selected with a correlation of 0,7 or higher with at least one of the other 29 FRUs. This is a small section of the set of HDDs, indicating that different usage patterns can be identified within a commodity, making them not commodity specific. After indexing the usages, the resulting usage pattern is displayed in Figure 5-4. When we analyze the characteristic values of the 30 FRUs that together create the usage pattern, the first observation is that all 5 divisions are present. With respect to the characteristic brand, 6 out of 7 possible brands are represented. For both characteristics, the results support the conclusion that Brand and Division cannot be used as indicators for a specific usage pattern.

The characteristic LTB Month forecast shows a range of forecasted amount between 1 and 120 pieces per 4-week period. These LTB Month forecast values represent different points in the PLC, but

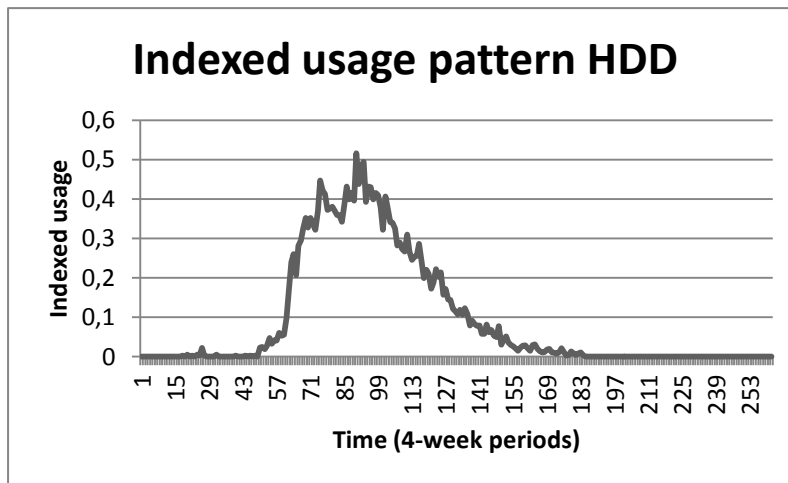


Figure 5-4: Indexed usage pattern based on correlated HDDs

are also based on different usage magnitudes between the FRUs. Because we combined the information of different FRUs to create the usage pattern, it is not possible to compare the results over the entire pattern. What we can conclude based on this selection, is that a single indexed usage pattern does not always result in a similar amount of items being used over time for different FRUs.

The final characteristic is Age FRU. The age range for the FRUs covers 1 to 11 years. There is a slight increase when the FRU is located more to the end of the PLC, but this is not the case for all FRUs. As a result, age as an indicator is not clearly related to a usage pattern or position in the PLC.

Based on these observations, we conclude there is no direct relationship between the usage pattern an FRU follows and the characteristic values it has. Therefore, characteristic values cannot be used to assess the usage pattern. And if we do not know the usage pattern, identifying in which stage the FRU is will be impossible. However, the added value of this information would also be limited, because it will not yield additional information about the usage pattern that will most likely be followed.

The only situation in which identifying the PLC stage the FRU is in will be possible, is when the usage pattern is known based on high correlation with an identified usage pattern based on complete data. Even when this is the case, there is still a high probability the FRU will not follow the correlated usage pattern exactly, because of actions outside the scope of control, like quality issues and substitutions. Another potential threat of limited historical usage data could be the placing the FRU at a different point in the PLC than is actually the case. But despite the uncertainty that will always be present, a known similar usage pattern might give guidance to what could be expected in the future.

5.7 Conclusion

Based on the results of the analysis performed in this section, the characteristics we identified as potentially influential to the usage pattern turn out to be unable to predict the partial usage pattern of an FRU. Therefore we conclude that characteristic values are FRU specific instead of usage pattern specific.

The only information that might give reliable information about the usage pattern is high correlation of the historical usage with the usage pattern of a (group of) FRU(s) for which the entire PLC is known. Because of data unavailability, it is not possible to test this. Data regarding historical usage from the beginning of the PLC of FRUs should be stored, to be able to do further analysis with respect to the existing usage patterns and the improvements that can be realized by incorporating these usage patterns in the forecasting processes.

6 What is the impact of the results on the current way of working?

In section 5 we concluded that usage patterns are not commodity specific, and FRUs cannot be clustered into groups with similar usage patterns easily and efficiently. These results raise questions about the applicability of CBLF. The basic assumption of this method is a similar usage pattern for FRUs belonging to a specific combination of Brand and Commodity. However, the results of section 5 do not point in the direction of a similar usage pattern for this combination. Especially the correlated usage pattern created in section 5.6 shows that a usage pattern can be similar for FRUs in different brands and divisions. Based on this result, we expect it to be likely that similar usage patterns can also be found for FRUs in different commodities, and usage pattern can differ within a commodity. This would violate the basic assumption of CBLF. Therefore we will test what the appropriateness of the basic assumptions in CBLF is by determining the performance of CBLF, derive potential causes for inaccuracies and methods to improve the performance.

6.1 What is the performance difference between CBLF and standard decline?

As stated earlier, the results of section 5 have lead to questions about the assumptions used in CBLF. To determine whether these assumptions are actually violated, we will compare the results for CBLF and forecasting based on standard decline. If the usage pattern for a specific FRU resembles the commodity specific PLC, the forecast accuracy of CBLF should be higher than the accuracy for the standard decline. If the assumption is violated, the results will be the other way around. We will test this assumption for the entire dataset and for the slow moving FRUs, to test whether the assumption holds for all cases or just for a selection of FRUs.

6.1.1 What is the aggregated performance for both methods?

To be able to assess the performance of the CBLF method, the PLC curve for the standard CBLF approach is used, without any modification. These results can then be used in a later stage to assess whether the adaptations made to CBLF will lead to more accurate forecasts. The forecast accuracy will be measured in terms of the average MAPE and MAD over the entire RSP of the FRUs, according to equations (6) and (7) respectively. Here MAPE indicates the average size of the forecast error in percentages of the actual usage, while the MAD gives an indication of the size of the forecast error. This performance measure will be determined for a specific set of 203 FRUs. These FRUs all belong to the division xSeries and have an LTB calculation executed in the year 2008. The division xSeries is selected, since for this division CBLF curves are available, making sure the PLCs used in the analysis are the same PLCs as those included into the pilot. We selected LTBs from the year 2008 for two reasons. The first reason is, that in the year 2008 only the standard decline (as explained in section 3.1.3) was applied in LTB forecasting. This will provide the certainty that we actually compare two different methods. The second reason is the applicability of both historical and future usage data, taking the LTB moment as the base time point. If we select FRUs with LTBs in 2008, there is at least 6 months of historical usage available, which can be used for the determination of the model fit. At the same time, there is a maximum of 4,3 years of data which can be used for testing the performance of the selected method.

When we analyze the results, for 108 FRUs the total amount forecasted based on standard decline is on average 139% more accurate than the total amount forecasted based on CBLF. For 56 FRUs, the CBLF results are 46% closer than the results with standard decline and 39 FRUs show no difference between the amount forecasted with CBLF and standard decline. These 39 FRUs are either FRUs that

have no usage in the time since the LTB, or the forecasts based on CBLF and standard decline are the same when rounded to an integer amount of pieces. The cases in which the forecasts are the same are all cases with low usage, since the actual usage over the 4 year period does not exceed 10 pieces per year. If we look at FRUs with an average usage of less than 10 pieces a year and an LTB Month forecast of less than 4, the difference in forecasted amount for both methods ranges between 0 and 10 pieces over the 4 year horizon. Most differences concern just 1 or 2 pieces, which indicates that the effect of applying one of these two forecasting methods is limited for slow moving FRUs. This is most likely due to the fact that both methods are not designed for slow moving FRUs.

	Number of FRUs	Forecast method	Aggregated average forecast accuracy	Aggregated median forecast accuracy
FRUs with usage in RSP	154	CBLF	MAPE: 314%	MAPE: 95%
		Standard decline	MAPE: 235%	MAPE: 80%
FRUs without usage in RSP	49	CBLF	MAD: 4,31	MAD: 0
		Standard decline	MAD: 4,08	MAD: 0

Table 6-1: Aggregated forecast accuracy CBLF and Standard decline

When we look at the MAPE results in Table 6-1, the difference between the average and median MAPE for the 154 FRUs that had usage in the time window using CBLF indicates there are a number of FRUs with extreme values. Although the median MAPE is substantially lower than the average, it is still a high deviation, indicating there is room for improvement in CBLF. For standard decline, the MAPE values point in the same direction as for CBLF, although performance is slightly better on an aggregated level. For the 49 FRUs that had no usage in the time window, the median MAD is zero for both forecasting methods, indicating that the large average deviations are the result of less than half of the FRUs. Although the difference between the average MAD for both methods is relatively small, it indicates that on average the difference between the actual usage and the forecasted amount using standard decline should be smaller than the difference between the actual usage and the CBLF forecast. This supports the observation that for most FRUs in the selection, standard decline is closer to the actual usage than the CBLF forecast.

Most accurate forecast method for FRUs with usage in the RSP	Number of FRUs	Aggregated average MAPE	Aggregated median MAPE
CBLF	49	203%	53%
Standard decline	95	241%	76%

Table 6-2: Aggregated forecast accuracy most accurate forecasting method

The overall performance indicates there are large percentage deviations between the forecasted and actual usage. If we solely focus on the method that provides the best results for the FRU, as listed in Table 6-2, we hope to see lower percentage differences, indicating a more accurate forecast on FRU level. For the 49 FRUs with usage for which CBLF provides more accurate results, the deviation between the average and median MAPE is lower than the average for all FRUs listed in Table 6-1, but it still means an inaccurate forecast according to Chien et al. (2010). For the 95 FRUs with usage for which standard decline provides more accurate forecasts, the average and median MAPE are higher than the CBLF values. This leads to the conclusion that if CBLF leads to the most accurate results, on average the percentage error is smaller. However, all values are above 50% deviation, and therefore provide room for improvement.

Concluding we can state that CBLF results in a more accurate forecast for 28% of the selected set, in 53% of the cases standard decline is more accurate and in the remaining 19% of the cases there is no difference between the performance of both methods. If we focus on the MAPE values, we can state that on an aggregated level standard decline leads to more accurate results. In the cases where CBLF is the most accurate method, the forecast accuracy is better than when standard decline is the most accurate. But since both the median and the average MAPE values for CBLF and standard decline are over 50%, both can be classified as inaccurate forecasting methods in general.

6.1.2 What is the performance for slow moving FRUs?

In literature, a substantial amount of papers can be found in which a method is developed especially for slow moving items. At SPO, the same method is applied to both fast and slow movers, which is contradictory to what one would expect based on literature. In this section we will determine whether the selected methods can be applied to both groups, or if performance is better for a specific group of FRUs.

The amount of slow movers in both methods will be determined, based on an LTB Month forecast of less than 4 and a maximum of 1 period historical usage as described in section 5.5, to see if there is a difference in the most accurate method between slow and fast movers. In Table 6-3 the amount of slow moving FRUs from the selection of FRUs forecasted most accurately with a specific forecasting method is listed. For the slow moving FRUs with usage in the RSP, the average and median MAPE values are calculated, to determine what the accuracy of the forecast with the specific forecasting method is. This value will be more positive than the aggregated value over all FRUs, since we only use a selection of most accurately forecasted FRUs. For FRUs without usage in the RSP, the aggregated MAD values are calculated.

Most accurate forecasting method	Amount of FRUs	Amount of FRUs indicated as slow mover	Usage in RSP?	Amount of slow movers	Aggregated average forecast accuracy	Aggregated median forecast accuracy
Standard decline	108	49	Yes	37	MAPE: 158%	MAPE: 89%
			No	12	MAD: 0,23	MAD: 0,08
CBLF	56	23	Yes	15	MAPE: 471%	MAPE: 59%
			No	8	MAD: 0,12	MAD: 0,10
No difference	39	39	Yes	10	MAPE: 85%	MAPE: 100%
			No	29	MAD: 0,01	MAD: 0,00

Table 6-3: Forecast accuracy slow moving FRUs

Based on the results for slow moving FRUs, we can conclude that slow movers cannot be forecasted accurately using either CBLF or standard decline, because the MAPE values for both methods are above 50%. The MAD values for all methods are relatively low, but when usage is zero, even a relatively low MAD value can lead to additional investments that were not required. A method specifically developed for slow movers is expected to provide more accurate forecasts.

6.1.3 What is the performance for fast moving FRUs?

In the previous section we determined the performance for slow moving FRUs, and concluded that both CBLF and standard decline cannot be categorized as methods that provide accurate forecasts for slow moving FRUs. This does not automatically implicate that accurate results can be realized for

fast moving FRUs. Therefore we also assesses the forecast accuracy of CBLF and the standard decline for the selection of fast moving FRUs, of which the results can be found in Table 6-4.

Most accurate forecasting method	Amount of FRUs	Amount of FRUs indicated as fast mover	Aggregated average MAPE	Aggregated median MAPE
Standard decline	108	59	290%	69%
CBLF	56	33	85%	41%

Table 6-4: Forecast accuracy fast moving FRUs

Based on the results in Table 6-4, our conclusion that when CBLF leads to a more accurate forecast the MAPE values are lower than when standard decline leads to the most accurate forecast is confirmed. When we compare the results with those of the slow movers in Table 6-3, the average MAPE value for the group most accurately forecasted with standard decline increases, while all other values decrease. This indicates that both methods are more suitable for fast movers, although there are some FRUs for which large inaccuracies still occur.

6.1.4 Conclusion

In this section we focused on calculating the performance of both CBLF and standard decline. On an aggregated level, standard decline has a lower MAPE value, indicating that this method predicts the usage more accurate than CBLF. This is confirmed by the results for fast moving FRUs. When we look on FRU level, we see that CBLF, standard decline or both can lead to the most accurate forecasts for a selection of FRUs. For the cases in which CBLF provides better results, the median is 53%, which is lower than the median value for standard decline. This indicates that if CBLF is the best method, the forecasts are relatively close for a large section of FRUs. However, the average MAPE is over 200%, indicating that the forecasted amounts can also differ largely from the actual usage. For slow moving FRUs, both methods provide inaccurate results, since both the median and the average MAPE values are over 50%. For fast movers, better results are realized, but improvement is still possible.

6.2 What are possible causes for inaccurate CBLF performances?

In section 6.1 we concluded that CBLF can provide the most accurate results for a specific subset of the selection, but large inaccuracies for others. In this section we will focus on possible causes for these inaccuracies. Information about the causes of inaccuracies is vital in determining and testing improvement options that could lead to more accurate CBLF forecasts. In the investigation of potential causes, we will first select a subset of FRUs, to be able to do a thorough analysis.

6.2.1 Which subset is to be investigated?

The assumption for CBLF is that every FRU in a specific commodity follows a similar usage pattern, but in section 5 we already concluded this does not have to be the case. As a result, the CBLF method could give accurate results for one FRU in a commodity, but not necessary for all. It might also be the case that exceptional usage patterns have influenced the CBLF curve in such a way, that the result is not accurate at all. In this section we will focus on investigating two possible causes that could lead to inaccurate CBLF results. We will select a set of commodities that provide both accurate and inaccurate forecasts, and test whether the causes are present. The commodities under consideration are the commodities as defined in the CBLF approach, which means they are different from the commodities we used in the previous chapters. A short description of the CBLF commodities can be found in appendix I.

Ideally, the selection of commodities to investigate should consist out of commodities for which none of the FRUs is forecasted most accurate with CBLF, commodities for which both methods lead to the most accurate forecasts and commodities for which CBLF always leads to the most accurate forecast. This composition would give the opportunity to investigate what the differences are between the groups, and which causes are related to those differences. However, based on the results we conclude that there are no commodities for which CBLF always leads to the most accurate results, because the maximum percentage of FRUs most accurately forecasted by CBLF is 33%. As a result, the third group from which commodities are selected is the group in which CBLF provides the most accurate result in at least 30% of the cases. The commodities assigned to each group can be found in Table 6-5.

CBLF never most accurate	Accurate results both methods	CBLF most accurate for $\geq 30\%$
AC-DC converter	Custom function cards	CD/DVD drives
Active backplanes	Fans	Display/video cards
Bezels/doors/fillers	Keyboards	Hard disk drives – ass
Ethernet adapters	Mech ass and sub ass	System planar
I/O riser cards	Memory subs	Tape drives
Rechargeable battery	Micro assemblies	
	SCSI adapter	

Table 6-5: Most accurate forecasting method commodities

In the analysis of potential causes, we will select 2 commodities from each group. The selected commodities are the AC-DC converter, I/O riser cards, Fans, SCSI adapter, System planar and Tape drives. In total 19 FRUs belong to these commodities. Of these 19 FRUs, 13 are forecasted most accurate with standard decline and the remaining 6 FRUs are forecasted most accurate with CBLF.

6.2.2 What is the impact of a differences in time on the market on the performance of CBLF?

One of the assumptions of CBLF is that every FRU in a certain commodity has a similar usage pattern over a PLC with similar length, in which the PLC length is determined by the point in time in which the smoothed CBLF curve reaches zero. Due to differences in the time the products are on the market, the created patterns might not represent the actual usage pattern accurately. For example, when the expected time on the market of the commodity based on CBLF is 12 years, while the actual time on the market of an FRU is 6 years, CBLF can lead to wrong forecasts, because the life cycle stages will not be located at similar points in time.

To determine whether the difference in time on the market has effect on the performance of CBLF, we analyze the differences of the 6 selected commodities. Based on the length of the CBLF curves, the expected time on the market for the FRUs in the commodities can be determined. The actual time on the market of the FRUs belonging to these commodities is determined by the time between the GA date and the EOS date.

When analyzing the performance of both forecast methods, the best analysis could be made after the FRUs are EOS, since all usage would have taken place by then, and the final difference between the forecasted and actual usage can be determined. In this selection, 2 of the FRUs are already EOS, the remaining FRUs not yet. For these FRUs we assume that if the FRU actually follows the CBLF curve, the actual usage over a section of the forecasted period should be similar to the forecasted amount.

Based on the differences in expected and actual time on the market, we can conclude that for the commodities AC-DC converter, SCSI adapter and System planer the expected time on the market is always longer than the actual time on the market. For these commodities, changes in technology that lead to obsolescence at an earlier point in time are the most likely cause of the difference. The other commodities show a mixed result, with both longer and shorter times on the market. With respect to the forecasted amounts, CBLF over forecasts the usage for 13 FRUs. With one exception, all these FRUs have an actual usage of less than 200 pieces over the 4 year time horizon. 5 FRUs are under forecasted, which all have a total usage of more than 450 pieces over the 4 year horizon. From this we can conclude that CBLF tends to over forecast low usage, and under forecast high usage.

For 11 FRUs the actual time on the market is shorter than the expected time on the market. If this is the case, CBLF is expected to over forecast the usage, since the curve will be cut off at an earlier point in time than it would end normally. If we analyze the forecasted results for these FRUs, the CBLF forecasts are higher than the actual usage for 7 FRUs. For the other FRUs the forecast is lower, which indicates that the usage for these FRUs was higher than expected based on CBLF. The under forecasted FRUs belong to 2 different commodities, consisting of both over and under forecasted FRUs. As a result, we cannot state that a shorter actual time on the market than expected always leads to over forecasting, nor that the CBLF curves for the 2 commodities with under forecasted FRUs are structurally under forecasting the usage.

For FRUs with the same actual as expected time on the market, it is expected the CBLF forecast should be close to the actual usage if the FRU follows the CBLF curve. There is 1 FRU in the selection for which the times on the market are the same. If we look at the forecasted results for this FRU, the CBLF forecast lies 80% higher than the actual usage. Since the actual usage is lower, we conclude that the FRU cannot follow the CBLF curve and equal times on the market do not automatically lead to accurate forecasts.

The remaining 7 FRUs have an actual time on the market which is longer than expected according to CBLF. If the FRU follows the CBLF curve, no usage will be forecasted for the years after the expected time on the market has passed. In theory, this should result in a CBLF forecast which is lower than the actual usage due to the years for which no forecast will be created. When we analyze the forecasts, we can see that for 5 of the FRUs the CBLF forecast is higher than the actual usage. This could be a positive effect in the case that the forecast turns out to be accurate at the end of the lifecycle, but it could also be an indication of the CBLF curves being inaccurate. Because the over forecasted FRUs are divided over 3 commodities with diffuse results, we again cannot state that the CBLF curve is inaccurate.

Based on these results, we cannot conclude that the CBLF approach leads to more accurate results if there is a difference between the expected and actual time on the market or when the times on the market are similar. This is due to the fact that the situations in which the approaches lead to better results differ and cannot be related to one of the options with respect to time on the market differences. What we can state is that the expected times on the market based on CBLF are inaccurate for most FRUs under consideration. As a result, more accurate estimations of the time on the market could result in more accurate forecasts.

6.2.3 What is the impact of an incorrect usage start date on the performance of CBLF?

In the current CBLF method, the historical usage is reallocated based on the GA date of the earliest predecessor, representing the start date of historical usage. A difficulty in this approach is, that using the GA date does not automatically implicate that the usage start date is correctly chosen. For some FRUs, usage starts already before GA, while for others usage might start at a later point in time. For the correctness of CBLF, correct reallocation is important though. When we do not use the correct usage start date, this can lead to large errors in the forecast. This is because the CBLF curves have a fixed period of increasing and declining usage, causing the effect that the reallocated position of the usage might not always correspond to the periods of increasing and declining usage. To test whether the reallocation ensures an accurate resemblance of the points in time in which usage increases or decreases, we will determine whether the forecast improves if the curves are placed at a different point in time, which we will further refer to as shifted.

To determine whether the usage is reallocated to the correct point in time we will focus on the forecasts created if we solely consider the GA date of the FRU, without looking at the predecessors. When we do not take into account the predecessors, the FRU will be shifted to an earlier point in the PLC then if the GA date of the predecessor will be used. Although not taking into account the GA date of the predecessors will not always lead to the best possible reallocation, we can conclude that if a different reallocation leads to a more accurate forecast, the probability the usage is reallocated correctly using the current approach is limited, and improvements can be realized by optimizing this approach.

Of the 19 FRUs under consideration, 8 have a predecessor. If we compute the forecast values while we do not take into account the GA date of the earliest predecessor, the CBLF forecast increases noticeably. For 6 out of the 9 FRUs, this leads to a forecast that is less accurate than the forecast with taking into account predecessors. For 3 FRUs, the forecast without a predecessor leads to a more accurate result.

In all cases, the normal CBLF forecast is substantially lower than the actual usage, while the CBLF forecast without predecessors is always higher. If we investigate the reason for this effect, we can see that the usage patterns show a fast decrease in the first periods, followed by a slowly stabilizing usage. The continuing fast decline of CBLF leads to under forecasting when the usage is reallocated. If the predecessor is not taken into account, the CBLF curve starts earlier, which adds a number of

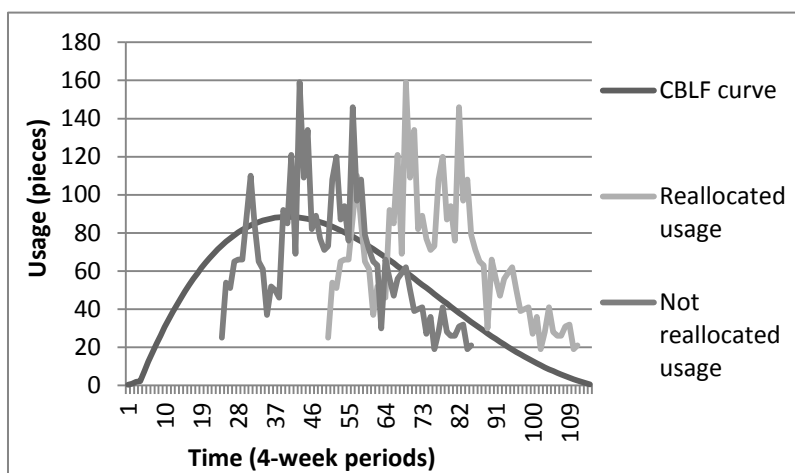


Figure 6-1: Example effect reallocating usage

periods with higher usage to the calculation. Although the results are mainly due to differences between the CBLF curve and the usage, we can conclude that not taking into account the GA date of the predecessor increases the forecast. This is especially positive for the 2 high usage FRUs, because for these FRUs CBLF tends to under forecast

usage with more than a thousand pieces.

Although we stated that not taking into account the GA date of the predecessor leads to improved forecasts for the 2 high usage FRUs, we do not know the reason why CBLF severely under forecasts the usage yet. If we analyze these FRUs, the CBLF curve and the actual usage seem to have a similar shape, but the start and peak moments of CBLF are located at an earlier point in time. Not taking into account the predecessor leads to a more accurate positioning in the curve and a more accurate forecast. This is illustrated in Figure 6-1, where the CBLF curve, reallocated usage and not reallocated usage for one of the high usage FRUs is shown.

Based on the results for the 8 FRUs with a predecessor, we can conclude that reallocation based on the GA date of the earliest predecessor does not always lead to reallocation to the correct point in time. At the same time, we cannot state that using the GA date of the FRU itself leads to correct reallocation, since actual usage might start at an earlier or later point in time than the usage in the CBLF curve. As a result, the reallocation of usage should be based on the usage pattern actually observed, to guarantee the best possible fit with the CBLF curve. This should also lead to improved forecasts.

6.2.4 Conclusion

On commodity level, different results can be realized with respect to the amount of FRUs correctly forecasted using CBLF, but in all cases the percentage of FRUs most accurately forecasted using CBLF is less than 30%. Main reasons for the inaccuracy of CBLF are differences in time on the market and a suboptimal reallocation of usage.

6.3 Which methods exist to improve the performance of CBLF?

In the previous section we concluded that the main reasons for the inaccurate performance of CBLF are differences in the time on the market and a suboptimal reallocation of usage. Another reason for inaccurate performances is a different usage pattern. In this section, we will focus on finding methods that minimize or remove the effect of these causes. As a result, the performance of CBLF should increase, and the method should be more flexible in adapting to the specific conditions of an FRU. This will result in the ability to make more accurate long term forecasts, which can be used in an LTB calculation.

6.3.1 Which methods can be applied to minimize the effect of differences in lifetime?

To overcome the issue of differences between the expected time on the market based on CBLF and the actual time on the market for a specific FRU, two different options can be applied. The first approach is to scale the CBLF curve, by determining the percentage of usage consumed per percentage of the time on the market passed. This approach provides the flexibility to adapt the CBLF curve to the actual time on the market of the specific FRU under consideration. In the remainder of this section we will refer to this approach as scaled CBLF. The second approach is to use standard distribution curves, for which the parameters will be determined based on historical usage of the specific FRU, to create the best possible fit. Standard distribution curves have the advantage that they will not reach zero usage, providing the opportunity to forecast usage for the entire time on the market. This is also the drawback of using standard distribution curves, because forecasts cannot be forced to zero at a certain point in time.

For the scaled CBLF curve, the percentage usage consumed per percentage of time on the market passed is determined by dividing the CBLF curve in 100 pieces with equal length, and determine the usage in a specific period based on the primitive of the smoothed CBLF function. When the usage in every period is known, the percentages of usage consumed in each period can be determined. The distribution of the percentage usage over the percentages of time on the market is fixed for a specific CBLF curve. Based on historical usage, the amount of pieces corresponding to 1% of usage can be determined for a specific FRU. When this information is known, combining the percentage of the time on the market passed and the amount of pieces representing 1% of usage, a forecast can be created for the remainder of the time on the market and/or for the percentages of usage remaining.

When we compare the forecasted amounts of CBLF and scaled CBLF for the 19 FRUs belonging to the selected commodities, 6 FRUs are more accurately forecasted using the scaled CBLF curve. These include all 3 FRUs of the commodity I/O riser card, indicating that the scaled CBLF curve is likely to provide more accurate results for this commodity. For the 3 other FRUs, the actual time on the market is longer than the expected time on the market, indicating that a longer forecast period might lead to more accurate forecasts for specific FRUs. Furthermore, for all 6 FRUs the percentage of time on the market passed is less than 30%, which confirms the importance of an accurate forecasting procedure early in the PLC.

Most accurate forecasting method	# FRUs	Average MAPE CBLF	Average MAPE scaled CBLF
Scaled CBLF	6	349%	98%
CBLF	12	193%	270%
Aggregated results	19	232%	202%

Table 6-6: Forecast accuracy CBLF and scaled CBLF

We calculated the MAPE values for the selection of 19 FRUs for both CBLF and scaled CBLF. The results of the 18 FRUs for which the forecasts differ are listed in Table 6-6. For 1 FRU, the forecasted amounts are the same for both methods. The results in Table 6-6 indicate that scaled CBLF can lead to forecasts that are more accurate than the forecast realized by CBLF, but also to more inaccurate results. However, on an aggregated level, the scaled CBLF approach is more accurate than CBLF.

Besides the scale issue, the current CBLF approach in which a polynomial is used to smooth the usage pattern creates the issue of curves that are cut off at a specific point in time to prevent negative usage. However, cutting of a curve does not indicate that there is no usage after that point in time, and a longer curve might lead to better results. A solution for this problem is the use of standard distribution curves. Some of the usage patterns and the CBLF curves resemble a shape that is similar to the shape of a standard distribution function, providing the opportunity to forecast the usage based on that distribution function for a longer time horizon, since the curve of a standard distribution function cannot drop below zero. Another advantage is that the number of parameters for a standard distribution curve is limited, and can be estimated for every single FRU. This results in parameters that are the most accurate for the specific FRU under consideration, which provides the opportunity to adapt the curve to the observed usage pattern.

To test whether standard distributions could lead to accurate forecasts, we decided to create an Excel model to forecast usage based on standard distributions. The incorporated distributions, which are explained in detail in appendix J, are the Weibull, Gamma, Exponential, Lognormal and Beta distribution. Besides standard distributions, we also added a forecasting function based on linear

regression and constant forecasting function. Finally, we also include the Bass distribution and Croston's method for intermittent forecasting, based on separately forecasting usage intervals and usage size.

The parameters for the distributions are estimated using the Excel solver. The goal is to determine parameters that result in the best possible fit on the available historical data, since we assume the future forecasts will be most reliable when the model has a good fit on the historical data. The most common method to determine the parameter values is the Least Squares method, in which the squared difference between the forecasted and actual usage is minimized to have the best possible fit (Johnson & Faunt, 1992). In the current CBLF approach, the determination of the parameter values is based on the highest possible correlation coefficient value R over the historical usage information though. Maximizing the correlation coefficient to estimate the parameters of curves is proposed by Brovetto, Delunas, Maxia & Spano (1989) in the area of physics, for cases in which the scale factor on the vertical axis can be omitted. Although omitting the scale factor in our case might seem strange at first sight, this is actually already the case in CBLF, because the forecast is based on the percentage differences of equation (4). As a result of this approach, the scale factor will always be 1 for the period in which the LTB is executed, and the scale values for all other periods are indexed accordingly. The only scale factor remaining to influences the forecast is the LTB Month forecast, but this factor is not included in the process of determining the parameter values.

A drawback of the approach of Brovetto et al. (1989) is, that the approach has not yet been tested on a situation similar to ours, because of the focus in the physics area. But comparisons can be made between both cases, and applying the correlation coefficient as indicator for parameter estimation would limit the changes to the current CBLF approach. As a result, the decision is made to determine the parameter values based on a Least Squares method and based on the correlation coefficient value. The comparison between both methods has the advantage that it provides the opportunity to assess whether the correlation coefficient can actually be used in estimating parameter values, or that the Least Squares approach is actually better. And when the Least Squares approach actually leads to better results, changing the objective of maximizing the correlation coefficient value in the parameter estimation process of CBLF might turn out to be an improvement option.

As mentioned, two approaches will be used to determine what the best parameter values are. The first aims to maximize the correlation coefficient R , which is calculated according to the following equation:

$$R = \frac{n \sum_{i=1}^n X_i Y_i - \sum_{i=1}^n X_i \sum_{i=1}^n Y_i}{\sqrt{n \sum_{i=1}^n X_i^2 - (\sum_{i=1}^n x)^2} \sqrt{n \sum_{i=1}^n Y_i^2 - (\sum_{i=1}^n Y_i)^2}} \quad (9)$$

Where X represents the historical data and Y represents the values of the created graph. The closer this value is to 1, the more accurate the created graph represents the historical data. The Least Squares methods we selected aims to minimizing the MSE, in which the MSE is calculated according to the following equation:

$$MSE = \frac{1}{n} \sum_{t=1}^n (x_t - \hat{x}_{t-1,t})^2 \quad (10)$$

Where n represents the number of periods in the time horizon, x_t the observed usage in period t and $\hat{x}_{t-1,t}$ represents the forecast made in period $t-1$ (or earlier) for period t .

We used the Excel model to calculate the forecasted amounts based on the parameters estimated by both methods according to the selected distributions for the 19 FRUs. The results show that for 12 FRUs one of the selected standard distribution gives a more accurate forecast than CBLF, for which the details can be found in Table 6-7, while for 1 FRUs the results are the same.

Parameter estimation procedure	# FRUs most accurately forecasted with one of the distributions
Correlation coefficient R	6
MSE	3
Same result for both	3

Table 6-7: Most accurate forecasts per parameter estimation method

The standard distribution leading to the most accurate forecast is not always the same, since these results are realized by the linear approach, Weibull, Gamma, Lognormal, Croston's, Constant and the Beta distribution. For the Lognormal distribution, promising results are realized when the parameters are estimated using MSE. For the Beta approach, the results are better when the parameters are estimated based on the correlation coefficient. For the other distributions, both parameter estimation methods can lead to accurate forecasts. This division is considered to be a logical result, due to the large amount of partial usage patterns available.

If we calculate the MAPE values on an aggregated level, clear differences can be seen between the forecast accuracy for the 19 FRUs based on the two parameter estimation procedures in Table 6-8. With the exception of the Exponential and Linear approach, a clear difference in favor of one of the two methods can be identified. An interesting observation is that the results do not always correspond with the individual forecast results. For the Lognormal distribution e.g., we stated that for the case in which the Lognormal distribution has resulted in the most accurate forecast, the parameters were estimated based on MSE. However, the average aggregated MAPE value is lower when the parameters are estimated based on R. Although the results indicate that both approaches could be used to estimate the parameter values, the limited size of the selection make it difficult to conclude that these results will be the same for a larger selection. More research will be required, but for the remainder of the research we select the parameter estimation method that leads to the most accurate forecasts for a specific standard distribution on an aggregated level.

Forecasting method	Aggregated average MAPE based on MSE	Aggregated average MAPE based on R
Weibull	63%	149%
Gamma	104%	71%
Exponential	90%	89%
Lognormal	550%	434%
Beta	283%	350%
Croston's	114%	137%
Linear	743%	743%
Constant	321%	258%
Bass	>10.000%	>10.000%

Table 6-8: Forecast accuracy standard distributions per parameter estimation method

Besides comparing the results of the forecasts created based on the two parameter estimation methods, we also compare the results to the forecasts made by CBLF. The average aggregated MAPE value is 379% for CBLF, which is substantially higher than the average aggregated MAPE values for the Weibull, Gamma, Exponential and Croston's approach.

Although the average MAPE value indicates that using standard distributions can lead to reasonably accurate forecasts, the probability that the correct distribution is selected for all FRUs under consideration is limited. When this approach is actually to be implemented, an approach needs to be determined based on which the most appropriate standard distribution and parameter estimation method are selected. This result supports the conclusion that the usage patterns differ within a commodity, because different distributions lead to the most accurate forecast within our review selection.

Concluding we can state that both scaling the CBLF curve as using standard distributions could result in more accurate forecasts, but the best method to select differs per FRU.

6.3.2 Which methods can be applied to minimize the effect of suboptimal reallocation?

In section 6.2.3 we discussed the concept of suboptimal reallocation, in which the moment in time the usage starts is not equal to the moment the usage is expected to start. As a result, the actual and expected usage curves are not running parallel. To minimize the effect of suboptimal reallocation, a suggested approach is to calculate the value of the correlation coefficient R according to equation (9) for all possible reallocation options and select the option with the highest value of R . This approach will be referred to as Shifted CBLF and is considered to provide accurate forecasts, because the usage is reallocated such that the usage pattern best follows the CBLF pattern.

Most accurate forecasting method	Number of FRUs	Average MAPE CBLF	Average MAPE shifted
CBLF	6	308%	474%
Shifted CBLF	13	199%	128%
Aggregated results	19	232%	237%

Table 6-9: Forecast accuracy CBLF and shifted CBLF

After the computation, for all but 2 FRUs the value of the correlation coefficient R improves when the start period is placed at a different point in time. For 14 out of the 19 FRUs, the difference in periods is less than 25, which is comparable with a change in position of less than 2 years. For 13 FRUs, the forecast after the LTB improves, while for 6 FRUs the shifted CBLF approach leads to substantial differences in forecasted amounts and results with CBLF are more accurate. The aggregated average MAPE results can be found in Table 6-9. From these results we can conclude that CBLF is more accurate in case the method leads to the best result for a specific FRU and slightly better on an overall level.

6.3.3 What methods can be used to optimize CBLF itself?

Besides solving the issues with suboptimal reallocation and suboptimal curves, improvement options also exist to optimize the basic principles of CBLF. At the moment, two different approaches can be identified regarding the creation of the CBLF curve, namely curve creation based on the average and on the summed indexed usage. These different approaches can lead to different results, as described in appendix C. Besides the different approaches, there is also no check to see if the usage patterns of the FRUs that are used to create the curve are actually similar. When this is not the case, the overall

pattern might not be as accurate as possible. Based on the FRU usage data used to create the CBLF curves for the commodity Tape drives, we will test what the effects of these issues are on the overall forecasts, and whether revising the approach would lead to more accurate forecasts.

Based on the original equations, the implemented CBLF curve is created using the sum of the indexed usages per period of time. In the version under construction, in which the usage of 2011 and the start

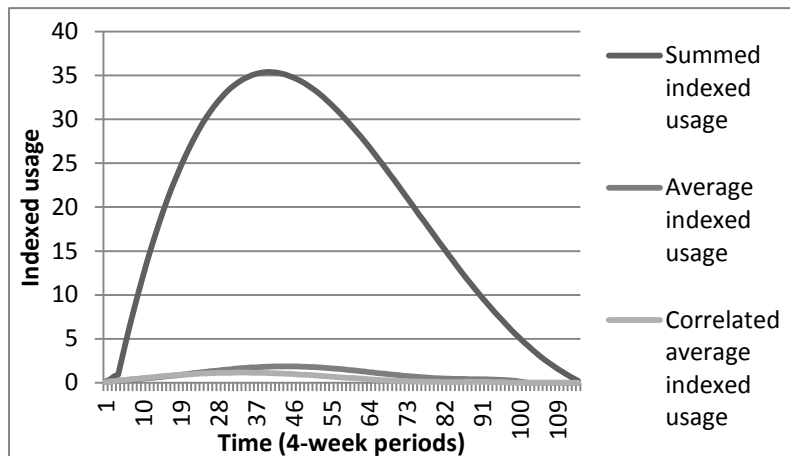


Figure 6-2: Difference summed and average indexed usage curves

of 2012 is added, the curve is created based on the average indexed usage per period of time. When we compare these curves, as shown in Figure 6-2, the curve based on the summed indexed usages is a lot steeper and slightly longer, which could result in forecast differences. Currently, 51 different FRUs are used to create the CBLF curve. For these 51 FRUs we create the correlation coefficient according to equation (9). When the correlation coefficient is larger than 0,6 we assume the usage patterns to be similar. In total, 18 of the 51 FRUs have a correlation of 0,6 or higher. This is only a limited fraction, indicating that different usage patterns are likely to exist within this commodity. When we create a curve for the 18 correlated FRUs based on the average indexed usage, the curve is slightly lower than the average indexed usage curve, as illustrated in Figure 6-2. These differences in results based on the same dataset proof that different curves can be realized easily, but this does not mean the curve is automatically correct. To determine which curve gives the most accurate representation of usage, testing is required.

For the 5 FRUs from the commodity Tape drives that have an LTB in 2008 we computed the results using the 3 different curves. The results show that all methods lead to the most accurate forecast at least once. The correlated average indexed usage curve tends to forecast intermittent usage relatively accurate, although this is not the pattern the curve is created for. FRUs that are already close to EOS are forecasted best with the average indexed usage curve, because the slower decline is a more accurate representation of the usage pattern. The FRUs that are at an earlier point in the lifetime are best forecasted with the summed CBLF curve, because of the steeper incline.

Forecast based on	Average MAPE when most accurate	Average MAPE when not most accurate	Aggregated average MAPE
Summed indexed usages	34%	63%	51%
Average indexed usages	1%	93%	74%
Correlated average indexed usages	32%	68%	54%

Table 6-10: Forecast accuracy different curve creation methods

Although the results in Table 6-10 are positive when applied to this small subsection of 5 FRUs, there is no guarantee that similar results will be realized for a larger subset of Tape drives. A larger subset would also provide the opportunity to test whether the position of the FRU in the PLC and the observed usage pattern can actually be related to the considered methods. For now we conclude

that all approaches can be considered as promising for a specific subset within the commodity Tape drives.

6.3.4 Conclusion

Different forecasting methods are investigated, based on reallocation, standard distributions, scaling, average indexed usage and correlated historical usage. For all options, positive results were realized when compared to CBLF for a selection of the FRUs, while for others results became worse. Based on these results we concluded that all options can lead to more accurate forecasts for specific FRUs, but also to larger inaccuracies for others.

6.4 Which of the possible methods is the most promising?

In this section, we will compare the aggregated forecasts over the RSP of the different methods and determine which method leads to the most accurate forecast and can therefore be labeled as the most promising approach. As a first step, we will combine the forecasting results for all methods to see which method leads to the most accurate forecast. The methods mentioned in section 6.3.3 are excluded, since we can only test these options on the commodity Tape drive, not for the other commodities, making comparing impossible. With respect to the results of the standard distributions, only the distributions that have most accurately predicted the usage of at least one of the FRUs are included. The forecasts made with the parameter estimation method that has the best aggregated results as listed in Table 6-8 are used in the comparison.

The best results are spread over different methods, and sometimes more than 1 forecasting method leads to good results. In Table 6-11 the number of times a specific forecasting method resulted in the most accurate forecast is listed, divided over slow and fast moving FRUs, where slow moving FRUs are labeled as FRUs having on average less than 10 pieces usage a year. When none of the slow or fast moving FRU is forecasted most accurately with a specific forecasting method, the entry n.a. will be inserted in the table, indicating we cannot calculate the results. The same entry is also used for the standard distributions, because the Weibull, Gamma and Bass distribution are not used to forecast slow movers, as are Croston's and Constant not used to forecast fast movers.

Forecasting method	Number of FRUs most accurately forecasted		Average MAPE when most accurate		Average MAPE when not most accurate		Aggregated MAPE results	
	Fast movers	Slow movers	Fast movers	Slow movers	Fast movers	Slow movers	Fast movers	Slow movers
Standard decline	0	2	n.a.	32%	133%	192%	133%	139%
CBLF	2	1	19%	0%	269%	283%	230%	236%
Scaled CBLF	1	3	10%	55%	280%	101%	259%	78%
Shifted CBLF	2	1	4%	14%	279%	276%	237%	233%
Weibull	1	0	6%	n.a.	67%	n.a.	63%	n.a.
Gamma	2	0	63%	n.a.	72%	n.a.	71%	n.a.
Croston's	0	1	n.a.	38%	n.a.	129%	n.a.	114%
Linear	2	1	1%	0%	906%	830%	767%	692%
Constant	2	0	6%	n.a.	390%	n.a.	331%	n.a.
Bass	1	0	12%	n.a.	44.907%	n.a.	41.454%	n.a.

Table 6-11: Number of good forecasts per forecasting method for a selection of 19 FRUs

There is a large variety in methods leading to good results for both the slow and fast movers, and the low amount of correct forecasts compared to the selection under investigation supports the

conclusion that there is not a single method that leads to good forecasts in all cases. The same conclusion applies to the commodities, since the preferred forecasting methods differ within a commodity.

To determine which of the method is the best overall, the MAPE value for the cases in which the method lead to the most accurate result, the cases in which the result was not the most accurate and the aggregated value over the 19 FRUs are calculated and listed in Table 6-11. Especially the aggregated MAPE value is considered to be an important indicator of the overall success when applying the method, because the probability of always selecting the correct method is low, and when this is the case, the deviation from the actual usage should be as small as possible.

Based on the results, we can conclude that all methods are accurate when they lead to the most accurate results compared to the aggregated performance. Based on the aggregated results, the Weibull distribution would be selected as the most appropriate method to forecast fast movers, while the scaled CBLF would be selected for slow movers. However, we only tested the results on a limited set of FRUs, which is not considered to be representative for the entire product range of IBM. As a result, we cannot state that one method is the best. Therefore we conclude that for fast movers, standard decline, Weibull and Gamma are the options that lead to the most accurate forecasts. For slow movers, standard decline, scaled CBLF and Croston's approach are the best options to consider.

6.5 Selection method

Although we have showed that different forecasting methods can result in more accurate forecasts covering a section of the RSP, we have not yet determined which information could be used to determine the forecasting method that should be selected at the LTB moment. For the SPO team to be able to identify which method would lead to the best results, indicators should be established to guide the team in their decision. We will try to determine these indicators for all forecast methods for the selection of 19 FRUs. We decided not to focus on the most promising forecasting methods determined in section 6.4, to be able to work with a larger section and to determine whether specific methods can be identified that always select the most accurate method.

Due to the differences in results with the different methods within every commodity, it is not possible to assign a standard method to the commodities. As a second option, the correlation coefficient R is considered. If the best method would have the highest value of R for all FRUs in the selection, it would be easy to select the method that leads to the best result. The problem in using the value of R as an indicator is, that there will not be any result when the forecast is constant over the entire time window, or when either the forecast or the actual usage is always zero. This will not be an issue in this case, because for cases with a constant forecast and usage the best approach can be determined easily without an indicator, and will therefore be excluded from our research. An important aspect to consider in the determination of the forecast method based on R is, that for the linear approach, high correlations are likely, since a declining section is quite common in the usage patterns, even though the overall usage pattern might not be declining.

First we select the forecasting method based on the highest value of R , categorized based on slow and fast movers, because we have already showed that different forecasting methods lead to good results for these types of FRUs. Therefore we expect there might also be a difference in selection method. When we calculate the average aggregated MAPE for the methods selected based on the highest R value, the results are less accurate than if we always would have selected the most

accurate forecast method, as can be seen in Table 6-12. If we look at the methods selected and the corresponding forecast accuracy, the largest errors occur when the linear approach is selected. To test whether the value of R for the linear method affects the result, we specified a selection method in which we select the method with the highest R value while excluding the linear R values. In this case, the average MAPE decreases for the fast movers.

Although the results of excluding the value of R for the linear approach are promising for the fast movers, the average MAPE value for the slow movers is still very high. Therefore we will use a second option, being the amount forecasted with the selected forecasting method over the period before the LTB of which historical usage is known. If this forecast is relatively accurate, it is expected that the forecast of the remaining PLC should also be accurate. When we solely focus on the forecasted amount, the aggregated performance improves for the slow movers. A striking observation is that the linear and constant methods lead to the most accurate forecasts for the historical usage, while this is never the case for the future periods. As a result, we exclude these methods for the second approach, which leads to a substantial decrease in the average MAPE value.

Approach	Number of FRUs with most accurate method		Aggregated average MAPE	
	Fast movers	Slow movers	Fast movers	Slow movers
Best forecast method	13	6	15%	47%
Best R	2	2	184%	207%
Best R with linear excluded	3	2	163%	207%
Best forecast PRE LTB	0	3	587%	122%
Best forecast PRE LTB with linear and constant excluded	0	4	387%	91%

Table 6-12: Performance forecasting approach selection methods

For both the fast movers and slow movers the average MAPE values are above 50%, which indicates that the results are inaccurate on average. However, both approaches can be used as an indicator of what usage might be possible. Whether the forecasts will improve is no guarantee, because it is not yet possible to assign the correct forecast method based on an indicator without further research. Excluding the linear and constant approach should also be seen as a conclusion based on this specific data set, the results might be different for a larger or different dataset.

6.6 Conclusion

On an aggregated level, standard decline leads to a more accurate forecast than CBLF, but when CBLF is more accurate, it is so substantially. Within both methods, room for improvement is available. For CBLF, improvement could be realized via optimization of the reallocation process, scaling the CBLF curve or by using standard distributions. We have tested these approaches on a subset of 19 FRUs, and all options have resulted in a more accurate forecast at least once. For fast movers, promising approaches are standard decline, Weibull and Gamma. Identification of the most promising method could be realized using the highest correlation coefficient value R while excluding the value for the linear approach. For slow movers, promising methods are standard decline, Croston's and scaled CBLF. The best identification method for slow movers is to use the most accurate forecast for historical usage over the period before the LTB, while excluding the value for the linear and constant approach.

7 How should IBM implement the results?

This research has focused on methods to cluster FRUs based on characteristic values, such that the usage pattern of the FRUs in one cluster are similar. Based on the research outcomes, we concluded that determining the specific usage pattern the FRU follows is not possible based on the currently available information, and that usage patterns appear to be FRU specific instead of commodity specific.

Based on this conclusion, we analyzed the performance of standard decline and of CBLF, and we proposed a set of improvement options for CBLF. These improvement options all provide accurate results for at least one FRU in the selection, again supporting the conclusion that the usage pattern is FRU specific. As a result, it is important to determine the best possible forecasting method for each FRU individually. Unfortunately, we did not succeed in finding an indicator that could be used to determine the best possible forecasting method for all FRUs, making it more difficult to implement the results.

To overcome the problem of a missing indicator, a solution has been developed. In this approach, an Excel tool will be created, which contains the forecasting options that show the most promising results on an aggregated level for both slow and fast movers. Historical usage data can be used as input, to determine whether the specific FRU is a slow mover or not, based on a boundary value of less than 10 pieces usage per year. The historical usage data can also be used to determine the best possible parameter settings for the combination of forecasting method and FRU, to maximize the probability of an accurate forecast.

The output of the Excel tool consists out of a graph with 4 possible usage patterns for fast movers, and 3 for slow movers. Furthermore, forecasts will be developed to cover the RSP, which should be inserted by the user. The possible usage patterns and corresponding forecasts can be used by the SPO team to visualize possible scenarios, and select the scenario that is considered to be most likely based on expert opinion.

On a general level, the tool should be seen as an intermediate step in the transition from point forecasting to range forecasting. Being able to forecast ranges accurately will provide the opportunity to determine the amount of risk SPO is willing to take in a well founded approach. However, at the moment range forecasting is not yet possible, because essential information is missing with respect to variability and the accuracy of the current forecasts. The Excel tool provides the option to get used to the idea of forecasting ranges, to learn which improvements might be realizable when the actual range forecasting method is developed and it could trigger the idea of urgency in researching and developing a full range forecasting model.

8 What are the effects when the results are implemented for all FRUs?

In our research we concluded that characteristics are FRU specific, and can only be used as an indicator for a set of potential partial usage patterns. Furthermore, we concluded that both the LTB forecasting approach based on standard decline and CBLF provide room for improvement, which could be realized via different methods. However, these methods also appear to be FRU specific, and we could not identify a straightforward method which always would identify the most accurate forecasting method based on the information available at the moment of the LTB. To be able to work with the results, the most promising methods are included in an Excel tool, that supports the LTB process by providing the opportunity to visualize possible usage patterns and their effect on the LTB need. In this section we will discuss the effects of the results on other FRUs and divisions, as well as the contribution the results can give to improving the forecasting process for all FRUs.

8.1 What is the applicability of the results to other FRUs?

Although the results neither lead to a straightforward clustering approach nor to an approach that assigns the most accurate forecast method to the FRU under consideration, the results can be applied to other FRUs.

We assume the possibilities for clustering based on characteristic values are limited and usage patterns differ per FRU for other FRUs as well. This is valuable knowledge in the forecasting process, because it triggers awareness and focus to determine the best possible approach for each FRU on a standalone basis. This should lead to more accurate forecasts, because we showed that it is possible to forecast future usage relatively accurate with one of the available methods. The difficulty will stay the selection of the best forecasting method, but using a combination of the Excel tool and the knowledge of the SPO team as a guideline to assess the best possible option within an, over time, developed framework should lead to more accurate forecasts.

The results also indicate that applying separate methods for slow and fast movers would lead to improved forecasts. Especially within divisions with a large percentage of slow movers, this might lead to substantial improvements over time.

8.2 What is the applicability of the results to other processes?

The performance determinations showed that CBLF could lead to accurate forecasts when applied correctly. In the day-to-day forecasting process, this knowledge can be applied to reduce the possibility of overstocking in the time before the LTB. A point of attention should be the monitoring of the correctness of the CBLF shape, and performing interventions like repositioning or scaling when required.

Besides the day-to-day forecasting, one of the considered forecasting options might also give valuable insights in the new products process, in which the amounts to stock initially need to be determined. Based on an initial forecast or on estimated market potential, the selected forecasting option might indicate how the usage could be spread over the PLC.

8.3 What are the advantages and disadvantages?

On an overall level, a set of advantages and disadvantages of the research outcomes and the Excel tool can be pointed out. The advantages are:

- The Excel tool provides the opportunity to consider the results of multiple methods in the determination of the LTB need, as an intermediate step towards range forecasting.
- The flexibility in the Excel tool provides the opportunity to adapt a forecasting method to that specific FRU.
- The limited set of identified partial usage pattern limits the amount of usage patterns to consider when analyzing the historical usage data of the FRU itself, which simplifies the process of assessing the potential future usage pattern.
- The forecast is created over 4-week time buckets, which ensures the method can easily be compared to the current forecasting approach.

The disadvantages are:

- It is not yet possible to determine the best forecasting method for each FRU without a large percentage of inaccurate decisions. Having the ability to correctly determine the forecast would lead to further improvement.
- Full usage patterns are not known yet, only partial usage patterns are identified.
- Uncertainty in usage patterns is not incorporated yet.

8.4 Conclusion

Implementing the Excel tool should lead to more accurate forecasts in the end, because the tool provides the SPO team the ability to visualize possible future usage patterns and corresponding LTB needs. This will not only be the case for the selection of FRUs we have incorporated in our research, but also on other FRUs, because the results are considered to be representative for the entire population.

Besides applying the results to the LTB forecasting process, additional information and visualization might also be valuable in the day-to-day forecasting process and in the new products process, to determine the amount of FRUs to acquire.

The flexibility and the additional insights the Excel tool can offer are considered to be the most important advantages, while the inability to select the most appropriate forecasting method is one of the largest disadvantages.

9 Conclusions, recommendations and further research

Based on the research, a number of conclusions and recommendations are determined and areas in which further research can lead to additional improvements will also be mentioned.

9.1 Conclusions

The most important conclusions that can be drawn from this research are:

- For approximately 70% of the FRUs overstocking already occurs in the process before the LTB calculation has to be executed.
- With respect to the accuracy of the LTB calculations, for most FRUs the required need is over forecasted, with on average 500%. For the slow and strong decline pattern, over forecasting is most common, with an average of 290%.
- For both day-to-day forecasting and the determination of the LTB need, the same methods are applied for slow and fast movers.
- 12 Different partial usage patterns are identified. These cannot be combined into complete usage patterns.
- Grouping FRUs on similar usage patterns based on partial usage data is not possible.
- Usage patterns appear to be FRU related and more likely determined based on actions outside the scope of control of SPO, like quality issues and the environment the machine operates in, than that they are commodity specific.
- Characteristics that might indicate the usage pattern are commodity specific, but value ranges of FRU characteristics that would indicate a specific usage pattern cannot be defined.
- Characteristic values, like the LTB Month forecast, could be used as an indication of a set of possible usage patterns.
- With the partial information currently available, the performance of the forecasts made with standard decline are often more accurate on an aggregated level than the forecasts made by CBLF.
- If CBLF results in more accurate forecasts, it is substantially closer to the actual usage.
- The main causes for inaccuracy of CBLF are differences between the expected and actual usage pattern, differences in the time on the market between the FRU and the CBLF curve, and a difference between the starting point of the actual usage and the expected starting point of the usage based on the curve used for forecasting future usage.
- For fast movers, standard decline, Weibull and Gamma are the options that lead to the most accurate forecasts. For slow movers, standard decline, scaled CBLF and Croston's approach are the best options to consider.
- Different forecasting methods lead to accurate forecasts within a set of FRUs, making the most accurate forecasting method FRU specific.

9.2 Recommendations

The following recommendations are based on the conclusions, in order of importance, could improve the processes further:

- Store historical usage data over a longer time period, until substantial information about complete usage patterns is available.

- Standardize commodity names, to prevent incorrect selections of commodities due to differences between CBLF and repair commodities.
- Forecast ranges instead of points when determining the LTB need, because of the variation in the usage per period between similar FRUs.
- Add performance measures and improve the day-to-day forecast by selecting the correct forecasting methods for slow and fast movers. Forecasting methods for slow movers that show good results in literature studies are e.g. single exponential smoothing and Croston's method, which are available in Xelus.
- Test the performance of forecasting methods for slow movers within an LTB situation. An initial test between Croston's method for intermittent usage and CBLF indicates that more accurate results might be realized by applying a method specifically designed for slow movers.
- Conduct further research to the improvement options that are incorporated in the Excel tool. Knowing which forecast method lead to the best result under certain circumstances would improve the knowledge and makes selecting a forecasting method easier.
- Use FRUs with high mutual correlations when creating the curve instead of using all FRUs from a specific commodity and to make sure the curve is always created according to the same approach, which should be clearly documented.
- A recommendation regarding the ease of use of the methods is to create a tool that makes sure the right predecessor and right GA date are selected, which makes the model easier to apply and the performance is secured, since the selection of a different GA date can have large impacts on the accuracy of the CBLF forecast.

9.3 Further research

Although different conclusions and recommendations can be based on this research, areas in which further research is valuable are defined, to be able to optimize the forecasting process even further. The areas of further research are:

- Research focused on the improvement potential of a forecasting methods that accurately forecasts slow movers, and the selection of the method that gives highest accuracy.
- Determination of the usage patterns available within IBM and the clustering and forecasting options over the full usage patterns, once usage data over the complete lifecycle of the FRUs are known.
- Further research based on full historical usage curves might clarify which method leads to accurate forecasts under specific circumstances.
- Investigate the effects of different forecasting approaches on a large sample that accurately resembles the populations of SPOs spare parts.
- Determine the variability in the usage and the accuracy of the LTB forecast, to make range forecasting possible.

References

- Bass, F. (2004). A New Product Growth for Model Consumer Durables. *Management Science* , 50 (12), 1825-1832.
- Bauer, H., & Fischer, M. (2000). Product life cycle patterns for pharmaceuticals and their impact on R&D profitability of late mover products. *International Business Review* (9), 703-725.
- Ben-Daya, M., Duffuaa, S., Raouf, A., Knezevic, J., & Ait-Kadi, D. (2009). *Handbook of Maintenance Management and Engineering*. London: Springer.
- Brovetto, P., Delunas, A., Maxia, V., & Spano, G. (1989). On Curve Fitting by the Utmost Correlation Method. *Il Nuovo Cimento* , 11 (11), 1645-1654.
- Chien, C.-F., Chen, Y.-J., & Peng, J.-T. (2010). Manufacturing intelligence for semiconductor demand forecast based on technology diffusion and product life cycle. *International Journal of Production Economics* (128), 496-509.
- Dinalog. (2010, 05 18). *Factsheet ProSeLo project*. Retrieved 02 10, 2012, from Dinalog: <http://www.dinalog.nl/files/publications/100518-factsheet-proselo-project.pdf>
- Dodge, Y. (2008). *The Concise Encyclopedia of Statistics*. Springer Science + Business Media, LLC.
- Eaves, A. (2002). *Forecasting for the ordering and stock-holding of consumable spare parts*. Lancaster: The Management School, Lancaster University.
- Huang, C.-Y., & Tzeng, G.-H. (2008). Multiple generation product life cycle prediction using a novel two-stage fuzzy piecewise regression analysis method. *Technological Forecasting & Social Change* (75), 12-31.
- IBM. (2011). *IBM Annual Report 2010*. New York: IBM.
- Johnson, M., & Faunt, L. (1992). Parameter Estimation by Least-Squares Methods. *Methods in Enzymology* (210), 37.
- Karmeshu, & Goswami, D. (2001). Stochastic evolution of innovation diffusion in heterogeneous groups: study of life cycle patterns. *IMA Journal of Management Mathematics* (12), 107-126.
- Kirshners, A., & Sukov, A. (2008). Rule induction for forecasting transition points in product life cycle data. *Scientific proceedings of Riga Technical University* , 170-177.
- Lapide, L. (2002). New developments in business forecasting. *The Journal of Business Forecasting* , 12-14.
- Meade, P., & Rabelo, L. (2004). The technology adoption life cycle attractor: Understanding the dynamics of high-tech markets. *Technological Forecasting & Social Change* (71), 667-684.
- Meenaghan, J., & O'Sullivan, P. (1986). The Shape and Length of the Product Life Cycle. *Irish Marketing Review* , 1, 83-102.

Meixell, M., & Wu, S. (2001). Scenario Analysis of Demands in a Technology Market Using Leading Indicators. *IEEE Transactions on Semiconductor Manufacturing* , 14 (1), 65-75.

NIST/Sematech. (2012, April). *e-Handbook of Statistical Methods*. Retrieved September 29, 2012, from Engineering Statistics Handbook: <http://www.itl.nist.gov/div898/handbook/>

Petrescu, E. (2009). A Statistical Distribution Useful In Product Life-Cycle Modeling. *Management & Marketing* , 4 (2), 165-170.

Rateq.com. (2012, April 30). *Historical Currency 20xx*. Retrieved May 1, 2012, from Rateq.com: <http://nl.rateq.com/historicalcurrency/2007>

Rink, D., & Swan, J. (1979). Product Life Cycle Research: A Literature Review. *Journal of Business Research* , 219-242.

Silver, E., Pyke, D., & Peterson, R. (1998). *Inventory Management and Production Planning and Scheduling*. John Wiley & Sons.

Singh, P., & Sandborn, P. (2006). Obsolescence driven design refresh planning for sustainment-dominated systems. *The Engineering Economist* , 51 (2), 115-139.

Solomon, R., Sandborn, P., & Pecht, M. (2000). Electronic Part Life Cycle Concepts and Obsolescence Forecasting. *IEEE* , 23 (4), 707-717.

Syntetos, A., Boylan, J., & Croston, J. (2006). On the categorization of demand patterns. *Journal of the Operational Research Society* (56), 495-503.

Wild, C., & Seber, G. (2000). *Chance Encounters - A First Course in Data Analysis and Interference* . New York: Wiley.

Wood, L. (1990). The End of the Product Life Cycle? Education Says Goodbye to an Old Friend. *Journal of Marketing Management* , 6 (2), 145-155.

Write A Writing. (n.d.). *Product Life Cycle (PLC): Stages, Development & Process*. Retrieved February 21, 2012, from Write A Writing: <http://www.writeawriting.com/business/product-life-cycle-plc-stages-development-process/>

Wu, S., Aytac, B., Berger, R., & Armbruster, C. (2006). Managing Short-Lifecycle Technology Products for Agere Systems. *Interfaces* (36), 234-247.

Appendices

- A. Divisions of IBM SPO
- B. The Product Life Cycle
- C. Example effect CBLF curve using average and summed indexed
- D. Usage pattern explanation
- E. Explanation commodities
- F. Results statistical tests
- G. Outlier analysis
- H. Box plots test selection
- I. Description commodities CBLF
- J. Standard distributions

A. Divisions of IBM SPO

The products IBM offers are divided over 7 divisions. These divisions are Lenovo, Mainframe, MVS, Power, Storage, System X and RSS. A short description of these divisions will be given.

Lenovo

Lenovo is the former laptop/desktop division of IBM. All activities in this division will be transferred to another company by the end of 2012.

Mainframe

A product from the Mainframe division (or System z) provides, within the product range of IBM, the highest level of availability and security. As a result, these systems are designed to meet the demand of the high-end market, and can be used for critical industries and applications, like banking. The capacity of the machines from the Mainframe division can be expanded, creating the opportunity to let the capacity grow with the business requirements.



Figure A-1: IBM Mainframe

MVS

MVS is the abbreviation of Multi Vendor Service parts, and represents all service parts IBM has in inventory that are not used for an IBM machine, but for machines of other vendors. IBM has these parts in inventory, because IBM has contracts in which they are required to service all machines of a customer, not just the IBM machines. Parts from different vendors are acquired and stocked to be able to service the non-IBM machines of the customer.

Power

Systems from the Power division consist of processors with high performance. Power systems are used to run and protect business applications, provide data storage capacity and backup options and encrypt data to protect files. These systems are used to keep business applications running and avoid unplanned system downtime.



Figure A-2: Examples of systems from the Power division



Figure A-3:
Disk storage
system on a
standalone
basis

Storage

Within the storage division, multiple storage products are combined, divided over tape systems and disk systems. All systems have different products in different price ranges and with different capacities. Storage products can be incorporated into other systems, or can function on a standalone basis. The systems can be used for day to day operations, but also for archiving e.g. data and emails.

Besides storing systems, data compressing systems are also part of the storage division. These systems can be used to compress real time data, in order to process more data in the same time.

System X

Servers in the X series are entry-level servers, focused on the low-end market, that can be used to run applications. To be able to run the applications, the system provides memory, storage and performance scalability. The X systems are easily deployed and managed and they offer the opportunity to expand memory, creating the flexibility to let the server grow with the company.



Figure A-4: System X processors

RSS

RSS stands for Retail Storage Solutions. This division delivers all kinds of sales hardware to retailers, like cash machines, self checkout terminals and receipt printers. The systems are available for different business sizes and different retail segments.

B. The Product Life Cycle at IBM

At IBM, every FRU has its own PLC, which shows the usage pattern of FRUs during the period of time the FRUs are used. Within IBM, the PLC is divided into five different phases, which are graphically represented in Figure B-1. The five different phases are:

- Pre General Announcement (GA)
- Early Life
- Mid Life
- End of Life (EOL)
- End of Service (EOS).

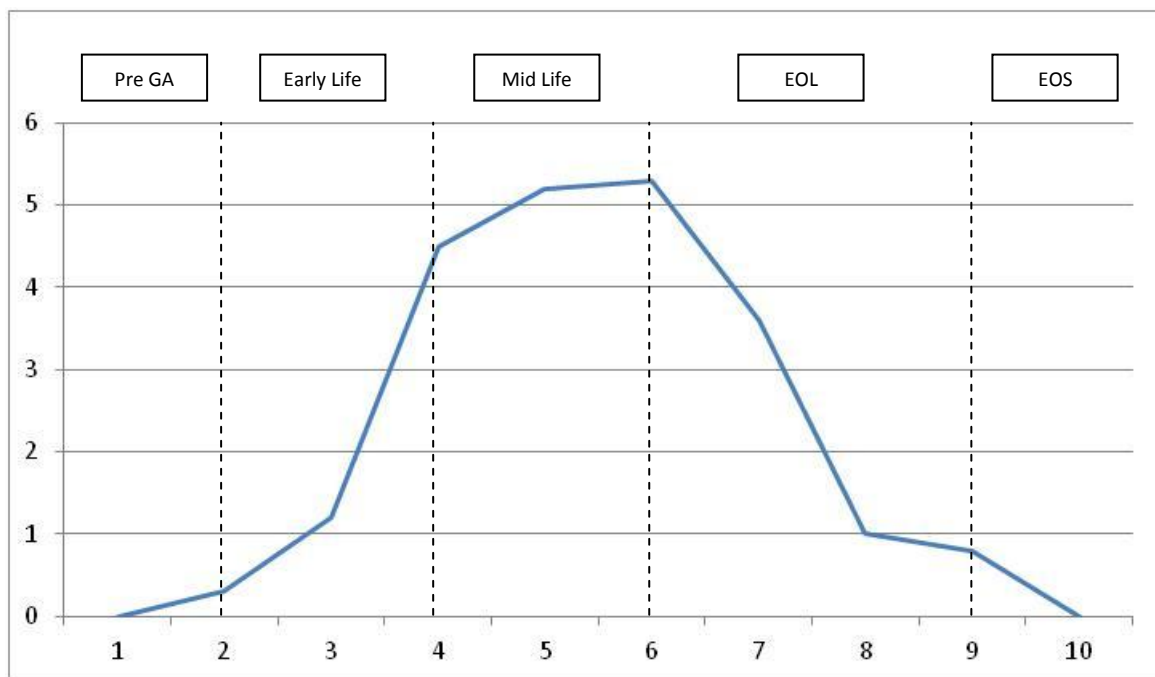


Figure B-1: Stages in a Product Life Cycle at IBM

In the first phase, the pre GA phase, the FRU is in the SPO system, but not yet available for the customer. The customer can place orders from GA date on. In this phase, SPO is responsible for acquiring and distributing the FRUs into the system, such that they will be available to the customer at the moment the GA date is reached.

The Early Life phase starts at GA, and will last until the processes of forecasting, ordering and distributing the FRU are working properly. During the Early Life phase, the FRU is managed by new product planners. These new product planners will estimate initial usage and the order quantities based on experience with similar FRUs.

In the Mid Life phase, the FRU is managed according to the standard forecasting procedure at IBM. The Mid Life phase continues until an End of Management and/or an End of Production notification is issued. When an End of Management and/or End of Production notification is issued, the supplier that issues the notification will not sell respectively not produce the FRU anymore. Other suppliers can still continue to produce and supply the FRU.

After an End of Management and/or End of Production notification is issued, the FRU reaches the End of Life phase. In this phase, the FRU is only used in the service department of IBM. This phase continues until the EOS date of the FRU is reached. In this phase, planning is based on the expected usage in the RSP. If a LTB has already occurred, the forecasting process stays the same, except for the fact that no new FRUs can be ordered from the supplier and SPO will actively monitor the FRU to make sure actions and adjustments to the parameters are applied to minimize the amount of FRUs left at EOS.

When the EOS date is reached, the EOS phase starts. At this phase, a decision needs to be made about the actions that need to be undertaken with respect to the FRUs that are still in inventory, but will not be used anymore due to EOS. Possible decisions are to transfer a number of FRUs to other countries that want to extend the service period or to scrap the remaining FRUs.

C. Example effect CBLF curve using average and summed indexed usage

In section 3.2 the method of CBLF is described. One of the issues of CBLF is the effect that summing the indexed usage for the FRUs can have on the calculated amount of FRUs to be ordered in the LTB compared to using average indexed usage. Here the described effect will be illustrated by an example, in which the PLC is determined for three parts that are considered to have a similar PLC with a length of 10 periods. In the first situation, we have partial information for parts 1 and 2, and for part 3 data is available for the entire PLC, as can be found in Table C-1.

Period	1	2	3	4	5	6	7	8	9	10
Part 1	1	3	4	5	6	8	6	5		
Part 2				5	7	9	4	3	2	1
Part 3	1	2	5	7	9	7	6	4	3	0

Table C-1: Partial usage data for example CBLF

From the usage data in Table C-1 we can determine the average usage per period according to equation (3). In this situation, the amount of periods in which usage could have occurred is 8, 7 and 10 periods respectively. As a result, the average usage per period for part 1 is 4,75 pieces, for part 2 it is 4,43 pieces and part 3 has an average usage per period of 4,4 pieces. These values will be used to determine the indexed usage according to equation (2). Subsequently, the sum of the indexed usage will be determined. The results can be found in Table C-2.

Period	1	2	3	4	5	6	7	8	9	10
Part 1	0,21	0,63	0,84	1,05	1,26	1,68	1,26	1,05		
Part 2				1,13	1,58	2,03	0,90	0,68	0,45	0,23
Part 3	0,23	0,45	1,14	1,59	2,05	1,59	1,36	0,91	0,68	0,00
Sum	0,44	1,08	1,98	3,77	4,89	5,30	3,52	2,64	1,13	0,23

Table C-2: Indexed partial usage data for example CBLF

In the second case, the information regarding the entire PLC is available for all three parts. This information can be found in Table C-3.

Period	1	2	3	4	5	6	7	8	9	10
Part 1	1	3	4	5	6	8	6	5	4	2
Part 2	1	1	3	5	7	9	4	3	2	1
Part 3	1	2	5	7	9	7	6	4	3	0

Table C-3: Full usage data for example CBLF

As for the previous calculation, the average usage per period can be calculated from the information in Table C-3, all based on the possibility that usage could have occurred in 10 periods. For part 1 the average usage per period is 4,4 pieces, for part 2 the average is 3,6 pieces per period and for part 3 it is on average 4,4 pieces. When the results are compared with the average usage per period for the partial data, the average is lower for part 1 and 2 in case all data is available. The resulting indexed usage, again according to equation (2), can be found in Table C-4.

Period	1	2	3	4	5	6	7	8	9	10
Part 1	0,23	0,68	0,91	1,14	1,36	1,82	1,36	1,14	0,91	0,45
Part 2	0,28	0,28	0,83	1,39	1,94	2,50	1,11	0,83	0,56	0,28
Part 3	0,23	0,45	1,14	1,59	2,05	1,59	1,36	0,91	0,68	0,00
Sum	0,74	1,41	2,88	4,12	5,35	5,91	3,83	2,88	2,15	0,73

Table C-4: Indexed full usage data for example CBLF

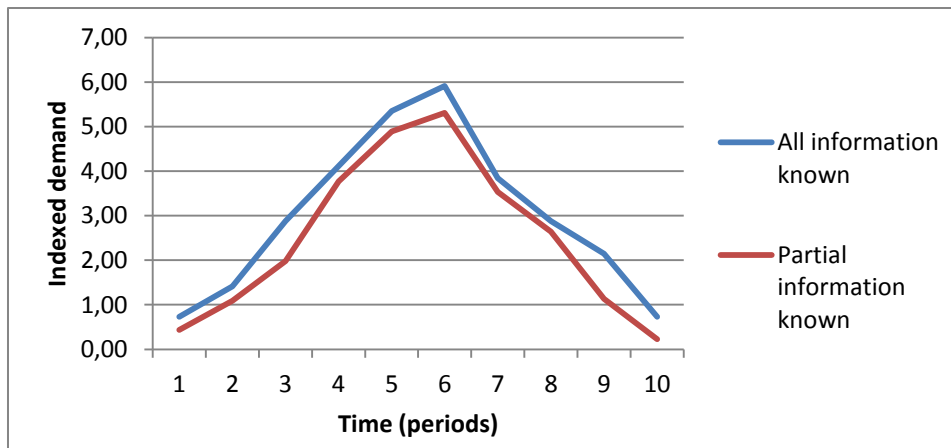


Figure C-1: PLC results for example CBLF

The resulting sum of the indexed usage for both situations are plotted in Figure C-1. There is a clear difference in the position of the graph on the y-axis, and a slightly different shape.

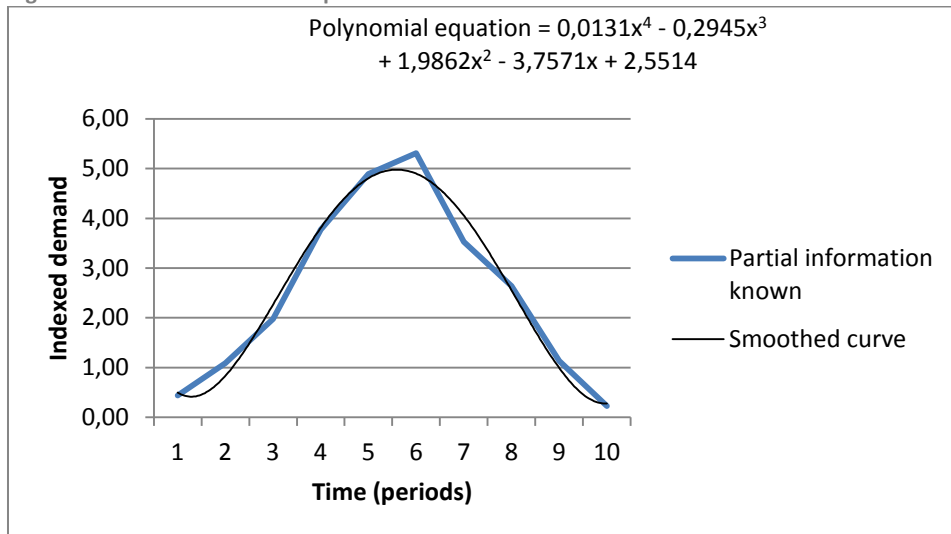


Figure C-2: Actual and smoothed PLC curve with partial information known

The next step is to smooth the graph, making it a more smooth PLC to work with, using a polynomial. In this case, a 4th degree polynomial is chosen for both situations, since there is little usage data available and the resulting equation follows the PLC quite accurate. The results of the polynomial and the corresponding equations for the situations in which partial and full information is known are shown in Figure C-2 and Figure C-3 respectively.

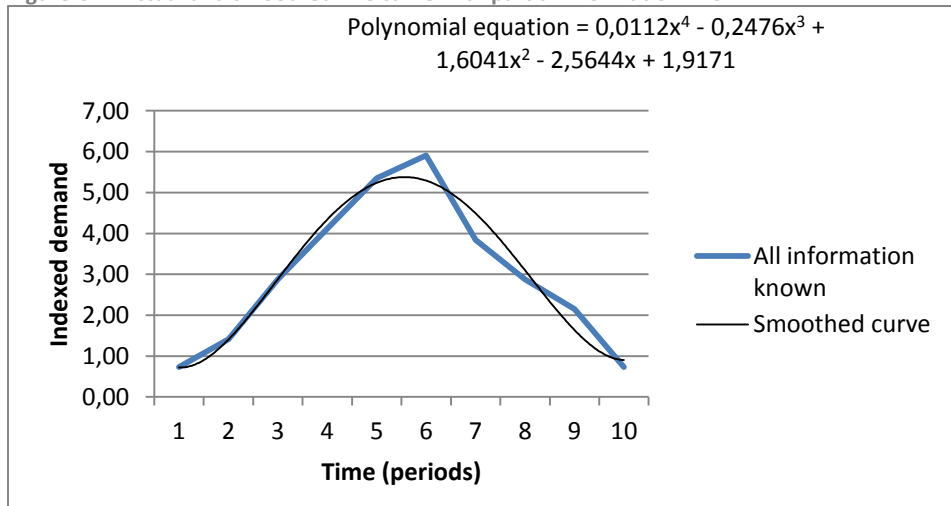


Figure C-3: Actual and smoothed PLC curve with full information known

In Table C-5 and Table C-6 the smoothed percentage differences are calculated for the partial and full usage situation, according to (4). Due to the differences in the indexed usage sum in the first period, the percentage differences largely differ over the life cycle.

Period	1	2	3	4	5	6	7	8	9	10
Equation value	0,505	0,846	2,282	3,829	4,823	4,910	4,053	2,528	0,927	0,154
Percentage difference	100%	168%	452%	759%	956%	973%	803%	501%	184%	31%

Table C-5: Percentages differences partial usage for example CBLF

Period	1	2	3	4	5	6	7	8	9	10
Sum	0,704	1,343	2,767	4,180	5,047	5,093	4,308	2,940	1,501	0,764
Percentage difference	100%	191%	393%	594%	717%	723%	612%	418%	213%	108%

Table C-6: Percentage differences full usage for example CBLF

To clarify the impact of differences in the percentage values, the expected requirement for the RSP will be calculated. This will be done for the situation in which the remaining expected requirement is calculated for the entire length of the PLC, with a usage of 5 pieces per period at the beginning of period 1. The results can be found in Table C-7.

Calculation	Results
Cumulative percentage difference partial usage (period 1 till 10)	4927%
Cumulative percentage difference full usage (period 1 till 10)	4068%
Expected requirement partial usage	$5 * 4927\% = 246,4$
Expected requirement full usage	$5 * 4068\% = 203,4$
Difference in expected requirement (partial – full usage)	43 pieces

Table C-7: Expected requirement results for example CBLF

From Table C-7 the conclusion can be drawn that in this example missing information results in an increased expected requirement of 43 pieces. For real situations, the difference in the resulting expected requirements can take on all kinds of values. In general, the amount of pieces to order with partial data is expected to be larger in situations with missing data at the beginning of the PLC and smaller when data is missing at the end of the PLC.

A potential approach to reduce this difference is to use the average indexed usage instead of the summed indexed usage. A short illustration of the effect of using average indexed usage will be given. In this example, we only compute the amount to acquire based on the average indexed usage, the smoothing step will be left out. It is expected that the effect of missing this step will be minimal, since smoothing is only executed to make the graph more easier to work with.

In Table C-8 the percentage differences, based on the average indexed usage and the summed indexed usage are given. The percentage differences are calculated according to (4), as described earlier. In Table C-9 the percentage differences for both situations can be found for the situation in which full usage was available.

Period	1	2	3	4	5	6	7	8	9	10
Avg. partial indexed usage	100%	248%	452%	574%	745%	808%	538%	402%	259%	52%
Sum partial indexed usage	100%	248%	452%	862%	1117%	1212%	806%	603%	259%	52%

Table C-8: Percentage differences partial usage

Period	1	2	3	4	5	6	7	8	9	10
Average full indexed usage	100%	193%	393%	562%	731%	807%	524%	393%	293%	100%
Summed full indexed usage	100%	193%	390%	558%	726%	810%	522%	393%	293%	103%

Table C-9: Percentage differences full usage

These percentage differences will again be used to compute the amount of parts to acquire. Here we also use the situation in which the remaining expected requirement is calculated for the entire length of the PLC, with a demand of 5 pieces per period at the beginning of period 1. The results can be found in Table C-10.

From Table C-10 we can conclude there is a large difference in the expected requirements. Especially for the case of the summed partial usage, the expected requirement is substantially larger, with approximately 80 pieces. Depending on the value of the FRU under consideration, this could lead to spending a substantially larger amount of money than required. In the current tendency of cost reduction, using the average indexed usage could therefore be an interesting improvement option.

Calculation	Results
Cumulative percentage difference avg. partial usage (period 1 - 10)	4178%
Cumulative percentage difference sum partial usage (period 1 - 10)	5711%
Cumulative percentage difference avg. full usage (period 1 - 10)	4096%
Cumulative percentage difference sum full usage (period 1 - 10)	4088%
Expected requirement average partial usage	$5 * 4178\% = 208,9$
Expected requirement sum partial usage	$5 * 5711\% = 285,6$
Expected requirement average full usage	$5 * 4096\% = 204,8$
Expected requirement sum full usage	$5 * 4088\% = 204,4$

Table C-10: Expected requirements average and summed indexed usage

Concluding we can state that in this example, using average indexed usage leads to a more accurate result, independent of the availability of all data.

D. Usage pattern explanation

In section 3.4.3 we identified 12 partial usage patterns. Here we will explain potential reasons behind the identified partial usage patterns.

No usage

An FRU can have no usage over a period due to several reasons. One of the reasons is that e.g. the FRU has a low probability of breaking down. This could be the case for several types of FRUs, but examples are cable or covers. These FRUs are important for the functioning of the machine, but can be used for a long time. Another reason could be that the FRU is just introduced and is not ordered yet, or the usage has stopped when the FRU is at the end of the lifecycle.

Single demand, 2-5 demands, 6-20 demands or 21-40 demands

The four patterns selected are all partial usage patterns in which usage is not continuous, but can occur in one period, followed by a number of periods in which there is no usage. This could for example be the case for FRUs that have a long lifetime, but break down once in a while. It could also be the case that the FRU belongs to a certain machine type of which the installed base is very small, which is automatically reflected in the usage pattern.

Relative stable

A relative stable usage pattern could e.g. occur when the FRUs have a relatively short period between breakdowns. The usage of these FRUs will then be relatively stable, since machines are sold at different points in time, and after a certain amount of time, all these FRUs need to be replaced. If the fluctuations in the amount of machines sold are limited, the usage will be relatively stable. Another possibility is that the machine is close to the end of the lifecycle, which will lead to a decrease in usage, making it more stable.

Increasing

An increasing usage pattern will most likely be the case shortly after the machine is introduced. At this point in time, the amount of machines being sold increases, which automatically affects the usage of the FRUs.

Bell shaped

In the bell shaped pattern, the usage first increases and then declines. This will most likely occur after the peak in the sales of the machines has passed, leading to a decline in the installed base. A lower installed base will result in less usage of the FRUs, because there are less machines that need maintenance.

Strong decline and slow decline

The strong decline and slow decline pattern will both appear at the end of the life cycle, when the amount of machines that require service decrease. As for the bell shaped pattern, less machines to service will result in less spare parts. Whether the pattern will be slow decline or strong decline could depend on several reasons, like the rate with which the installed base decreases, contract specifications, the lifetime of the part, et cetera.

Fluctuating

A fluctuating pattern could occur based on several reasons. A specific type of FRU could have for example a higher probability to break down in the summer, due to e.g. overheating. Another reason

could be that maintenance is referred to certain periods, which would also increase the usage in those periods. At a teaching institute e.g., it could occur that most maintenance is executed during the summer holiday. When a large amount of teaching institutes applies this practice, usage will increase in that period.

Extreme peak

An extreme peak could be due to two reasons. One of the reasons is that when one FRU breaks down, the other FRU(s) follow quickly, because they need to compensate for the broken one. This could happen for e.g. hard disks. Another option is related to environmental influences. A fire in an office building could destroy a large amount of FRUs, for which an order will be placed to replace them. This is an exceptional order, but will influence the usage pattern.

E. Explanation commodities

A commodity is a group of similar FRUs, that can be regarded as similar in the sense of their purpose. To give the reader a better understanding of the commodities, an explanation of the commodities from Table 3-5, with an LTB value of over €1 million, will be given here.

CARD

The first commodity to discuss is the commodity CARD. This commodity consists of all kind of cards that are required for the machine to work. An example of a card is a video card, which is required when the machine has a screen on which graphics are displayed.

CPU

A CPU is a processor, which is required for the functioning of the machine. There are different types of CPUs, which are all combined in the commodity CPU.

HDD

The commodity HDD consists out of Hard Disk Drives, which are required for storage. There are different types of HDDs, depending on the requirements of the machine.

MEM

The commodity MEM consists out of all kinds of products that provide additional memory, like memory cards.

MISC

MISC stands for miscellaneous, and consists out of all kinds of FRUs. Most of these FRUs are covers, screws and other FRUs that are used in machines, but do not perform a vital function.

PLNR

PLNR stands for Planar, which is the motherboard of a machine. This is the main system of the machine, which cannot function without it.

POWR

The commodity POWR consists out of power supplies, which provide the power required for the functioning of the machine.

TAPE

The TAPE commodity consists of all storage products that use tape as a medium to store data on. The principle is similar to the principle of music cassettes that was used before the CD was used for storage.

F. Results statistical tests

In section 5.2 we focused on determining whether there is a relationship between the characteristics under consideration and the potential usage patterns we have identified. As a basic assumption, we assume that the characteristic values are different for the 12 partial usage patterns, because this would indicate that a specific characteristic value can be assigned to a single partial usage pattern. Every time this characteristic value is observed for a specific FRU, the partial usage pattern would be known and can be applied in the forecasting process.

After determining the variable types and the corresponding statistical test that would lead to the most appropriate results, the first step is to test the correlation between the partial usage patterns and the characteristic values for the model building selection. All tests are executed in the program SPSS Statistics, with the partial usage patterns as independent variable and the characteristics as dependent variables.

For the characteristics Brand, Commodity, Division and TM144 Status the Chi-Squared test of Independence is applied. We will test the hypothesis that the results can be found without any relationship between the characteristics. This means that the division of the partial usage pattern over the characteristic values is random, and there is no direct relation between a characteristic value and a partial usage pattern. If we reject the hypothesis, this relationship is likely to exist. We will reject the hypothesis when the p-value is less than 0,05. This indicates a probability the results are found without any relationship between the variables is less than 5%, which is very small.

Characteristic	p-value Chi Square (2-sided)	Hypothesis accepted?
Brand	< 0,001	No
Commodity	< 0,001	No
Division	< 0,001	No
TM144 Status	< 0,001	No

Table F-1: Results Chi Squared test of Independence

In Table F-1 the results of the Chi Square test of Independence for the characteristics Brand, Commodity, Division and TM144 Status are listed. We clearly see that the p-values are less than 0,05 for all characteristics, indicating that we reject the hypothesis of a missing relationship between the variables. This supports the conclusion that these characteristics might influence the partial usage pattern of the FRU.

For the characteristics Age FRU, Forecasted reliability, LTB Month forecast and RSP we apply the Kruskal-Wallis approach. In this case, we will test the hypothesis that the average value for the characteristics is the same for all partial usage patterns. Although slightly different formulated, in principle we still test whether the results can be found without any relationship between the variable. As for the Chi Squared test of Independence, we again reject the hypothesis when the p-value is less than 0,05.

Characteristic	p-value Kruskal-Wallis	Hypothesis accepted?
Age FRU	< 0,001	No
Forecasted reliability	< 0,001	No
LTB Month forecast	< 0,001	No
RSP	< 0,001	No

Table F-2: Results Kruskal-Wallis

In Table F-2 the results for the Kruskal-Wallis test are listed. As was the case for the Chi Squared test of Independence, we can again clearly see that the p-values are less than 0,05. Therefore these characteristics can also be related to a partial usage pattern.

A large drawback of this approach is that we performed the test for a set of 1936 FRUs. There is a probability that there is a relationship between a part of the characteristic values and partial usage patterns, but not between all. As a result, some characteristics might be related to a partial usage pattern under specific circumstances. This is something we cannot assess based on the tests for the entire selection, because we now only know there is a relationship between the characteristics and the partial usage patterns.

To assess whether it is actually the case that certain characteristics are only related to the partial usage patterns under specific circumstances, we will perform statistical tests on a number of subsets, to see whether results are similar. Here, the subsets under consideration are the commodities. We have selected the characteristic Commodity because the results of the test would give us the ability to assess whether it is actually the case that specific characteristic values can be related to a specific partial usage pattern within a commodity. When this is the case, forecasting based on CBLF would become easier because the partial usage pattern, and therefore the stage the FRU is in, is known. The commodities we will focus on are the 8 commodities that have a total LTB value that exceeds €1 million. These are all large commodities, and being able to assign partial usage patterns to the FRUs in these commodities based on characteristic values provides a large potential for cost savings.

For the characteristics Brand, Division and TM144 Status we use Fisher's Exact test, because of the limited amount of observations. The hypothesis we will test is again the hypothesis that the results can be found without any relationship between the characteristics, which will be rejected when the p-value is less than 0,05. The characteristic Commodity is not incorporated in the test, because within a single commodity, the value cannot change, and the test results cannot lead to anything other than an acceptance of the hypothesis.

Because we can only perform Fisher's Exact test for two variables that each can take on two values, multiple tests are required to cover all 12 partial usage patterns and all possible characteristic values. This also results in different p-values, that either reject or accept the hypothesis for a specific combination of variables. In assessing whether the hypothesis will be accepted or rejected on an aggregated level over all 12 partial usage patterns, we assume that the hypothesis is rejected when in more than 50% of the characteristic value groups at least 1 combination shows a p-value lower than 0,05.

Commodity	eServer	ITS	RSS	Storage	xSeries
CARD	8	14	0	15	23
CPU	42	27	0	14	27
HDD	6	0	0	5	7
MEM	4	11	4	3	22
MISC	4	4	0	4	10
PLNR	5	12	1	6	10
POWR	9	5	0	5	1
TAPE	0	0	0	2	2

Table F-3: Number of cases with a p-value less than 0,05 for characteristic values for Division

For the characteristic Division, 312 combinations of characteristic values are required to cover the entire value range for both characteristics. In Table F-3 the results per division for the 8 commodities are listed. Based on our requirement that at least 50% of the characteristic values should have at least 1 case in which the p-value is less than 0,05 we can state that the hypothesis is rejected for the commodity TAPE. For the commodities MEM and PLNR, there appears to be a relationship between the characteristic values and the partial usage patterns. For the remaining commodities, the results indicate a positive result for a part of the divisions, which could indicate that the characteristics could only be related for a selection of the partial usage patterns.

Commodity	Disk products	Point of Sale	Power systems	Server	System i	System p	Tape products	zSeries
CARD	22	0	3	30	13	23	6	3
CPU	0	9	24	41	34	36	14	0
HDD	6	0	1	6	1	6	0	0
MEM	4	4	0	22	13	7	0	0
MISC	9	0	2	16	8	11	7	5
PLNR	6	3	2	10	3	12	0	0
POWR	0	0	0	0	2	2	0	0
TAPE	3	0	0	2	2	0	7	0

Table F-4: Number of cases with a p-value less than 0,05 for characteristic values for Brand

For the characteristic Brand, of which we can find the results in Table F-4, 546 possible combinations are required to cover all potential combinations. Based on the results, we conclude that the for the brands POWR and TAPE the results do not exceed the 50%, and therefore we conclude that there is no clear relationship between these brands and the partial usage patterns. For the other brands, the results point again in the direction of a relationship with specific subsections instead of a relationship with the entire characteristic value range.

Commodity	TM144 Status
CARD	3
CPU	0
HDD	1
MEM	1
MISC	7
PLNR	8
POWR	0
TAPE	1

Table F-5: Number of cases with a p-value less than 0,05 for characteristic values of TM144 Status

For the characteristic TM144 Status, there are only 2 possible characteristic values. As a result, 78 combinations cover all combinations of the partial usage patterns, and therefore the entire value range. In assessing whether or not we will reject the hypothesis, the initial prerequisite of more than 50% of the characteristic values having at least 1 case with a p-value less than 0,05 is realized as soon as the first observation is found. Therefore we change the prerequisite in more than 10% of the possible cases has a p-value less than 0,05. This indicates a boundary value of 8 observations, making the commodity PLNR the only commodity for which the TM144 Status might have a relationship with the partial usage patterns. The commodity MISC almost reaches the boundary value, indicating that this might also be an interesting commodity to take into account.

As a final step we will focus on the characteristics Age FRU, Forecasted reliability, LTB Month forecast and RSP. These characteristics are again tested with the Kruskal-Wallis approach, since this test can also be applied to small samples.

FRU Commodity	Age FRU	Forecasted reliability	LTB Month forecast	RSP
CARD	<i>0,036</i>	<i>0,007</i>	<i><0,001</i>	<i><0,001</i>
CPU	0,460	<i>0,001</i>	<i><0,001</i>	<i><0,001</i>
HDD	<i>0,005</i>	0,119	<i><0,001</i>	0,078
MEM	0,160	0,678	<i><0,001</i>	<i><0,001</i>
MISC	0,123	<i><0,001</i>	<i><0,001</i>	0,257
PLNR	<i><0,001</i>	0,586	<i><0,001</i>	0,001
POWR	<i>0,002</i>	0,414	<i><0,001</i>	<i>0,010</i>
TAPE	0,032	0,158	<i><0,001</i>	0,098

Table F-6: Resulting p-values hypothesis testing of characteristics on commodity level using Kruskal-Wallis

The resulting p-values of the Kruskal-Wallis test can be found in Table F-6, in which all values smaller than 0,05 are displayed in italic and are considered to be potentially influential. In Table F-6, certain p-values are substantially larger than the boundary value of 0,05 for a characteristic to be considered as influential. For the characteristics Age FRU, Forecasted reliability and RSP, there is a clear distinction between commodities for which the results point in the direction of a relationship between a partial usage pattern and a characteristic value, and cases in which this is not the case. The only characteristics for which the hypothesis is always rejected is the LTB Month forecast, making this the most interesting characteristic from the selection.

Because all test results were significant for the complete dataset and pointed in the same direction, we conclude that the characteristics that potentially influence the PLC depends on the commodity of the FRU, and might even hold for a specific subsection within a commodity. As a result, we expect that the value ranges the potentially influential characteristics take on will also differ per commodity.

G. Outlier analysis

Before we can test whether the characteristic value ranges can be used to identify partial usage patterns in section 5.3, the first step is to perform an outlier analysis. Outliers are values of a variable that are considered to be inconsistent when compared to the majority of the variable values. They can occur due to measurement errors, but the results can also be an anomaly. Including outliers in the selection could lead to incorrect observations, because their extreme values will affect e.g. the mean value. If these incorrect values will be used to draw conclusions on, the final conclusions will also be incorrect.

The outlier analysis will be based on the 3-sigma rule, in which the outliers are categorized as values outside the range $[\bar{x} + 3\sigma, \bar{x} - 3\sigma]$, with $\bar{x}$ being the mean and σ being the standard deviation. This approach is originally designed for variables that are normally distributed. Then 99,7% of the variables should fall within the range. For other distributions, at least 89% of the observations fall within this range (Wild & Seber, 2000).

Complete HDD model selection

We will perform an outlier analysis on the characteristics LTB Month forecast and Age FRU for the entire HDD range, to make sure the conclusions in section 5.3 are based on valid data.

The first range will be determined for the LTB Month forecast, which has a mean value of 20,79 and the standard deviation is 50,94. Based on these values, the range of allowable values becomes $[-132,16 ; 173,74]$. All values outside this range will be categorized as outliers. In the selection, two FRUs have an LTB Month forecast value larger than 173,74 and are excluded from the research. Both FRUs belong to the fluctuating partial usage pattern.

The second range will be computed for the characteristic Age FRU, which has a mean of 6,15 years and a standard deviation of 4,26 years. When we combine these values, we will get the allowable range of $[-6,63 ; 18,94]$. The largest age value in the selection of HDDs is 18 years, which stays within the boundaries. Therefore we conclude that there are no outliers for the variable Age FRU.

Increasing usage pattern

In section 5.4 we focus on combining characteristic value ranges to identify the increasing usage pattern. As for the complete HDD selection, we will also identify outliers for the increasing usage pattern, to make sure our conclusions regarding the possibilities of identification with a group of characteristics are not influenced.

For the LTB month forecast, the mean value is 16,99 with a standard deviation of 28,22. Combined this leads to a range of $[-67,67 ; 101,65]$. The largest value of the LTB forecast within the increasing usage pattern is 92,46. This value stays within the range, so we do not need to exclude any values. For the characteristic Age FRU, the mean is 2,45 and the standard deviation is 1,97. The allowable range is $[-3,45 ; 8,36]$, while the maximum age value is 6. As a result, we do not have any outliers for the increasing usage pattern.

Strong and slow declining usage patterns

To determine if the results from section 5.4 for the increasing usage pattern are representative for other partial usage patterns, we test the results for a combination of the strong and slow declining

patterns, because both patterns represent a similar shape. To do this, we also perform an outlier analysis for the combination of these patterns.

For the LTB Month forecast, the mean is 30,04 with a standard deviation of 34,77. The resulting range is [-74,26 ; 134,33]. The largest LTB Month forecast value is 117,35. This value is within the range, so there are no outliers in this section. For the characteristic Age FRU the mean is 6,96 and the standard deviation is 3,50. Combining, the range is [-3,55 ; 17,47]. With a largest Age value of 16 years, we do not identify any outliers.

H. Box plots test selection

In section 5.3 we created box plots to determine the differences in the value ranges for the characteristics under consideration. The box plots for the model selection are displayed in section 5.3, but the box plots for the characteristic value ranges of the test selection can be found here.

The test selection consists out of 47 HDDs. For these HDDs, the box plots are created using the same approach as for the modeling selection.

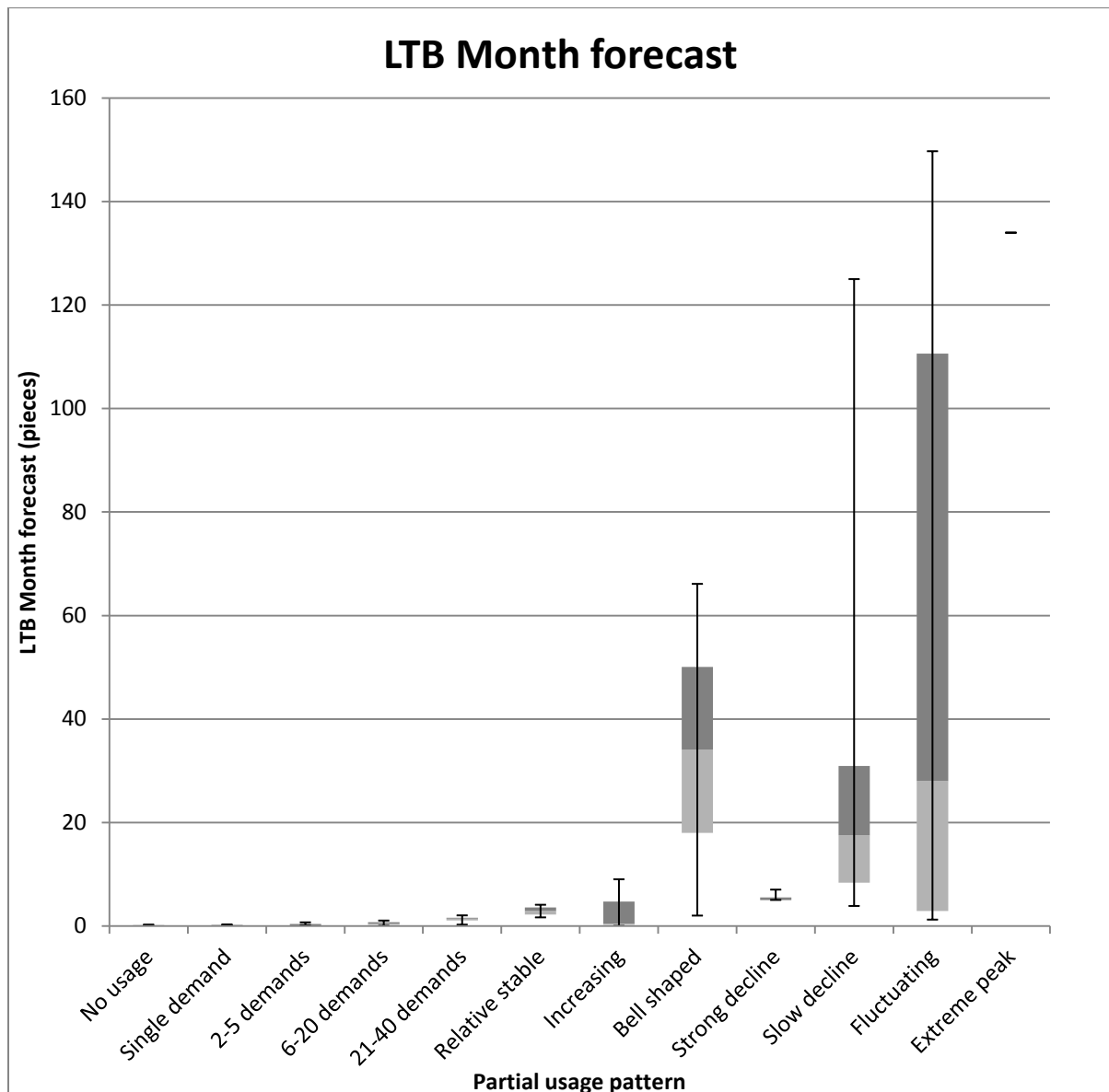


Figure H-1: Box plot LTB Month forecast test selection

In Figure H-1 the box plot of the LTB month forecast values is displayed. The slow decline, bell shaped and fluctuating patterns show a large range of forecast values used in the LTB calculation. Since the ranges for these partial usage patterns in the model selection were also large, this result is a confirmation of the large forecast range corresponding to these partial usage patterns. The other partial usage patterns all show a low forecast value, which makes it difficult to distinguish which patterns the FRU could follow based on these results. For the 5 most right partial usage patterns in

Figure H-1 this means a confirmation of the results of the modeling selection, but for the others the results are different.

Figure H-2 represents the box plot of the Age FRU for the test selection. If we analyze the results in this figure, the age ranges for all partial usage patterns are overlapping. Some are more concentrated on lower ages and others cover the entire range, but there is no clear difference in the age range covered for a specific partial usage patterns. Because similar results are found in the modeling selection, the results for the test selection confirm the conclusion that age is not a clear indicator for the partial usage pattern.

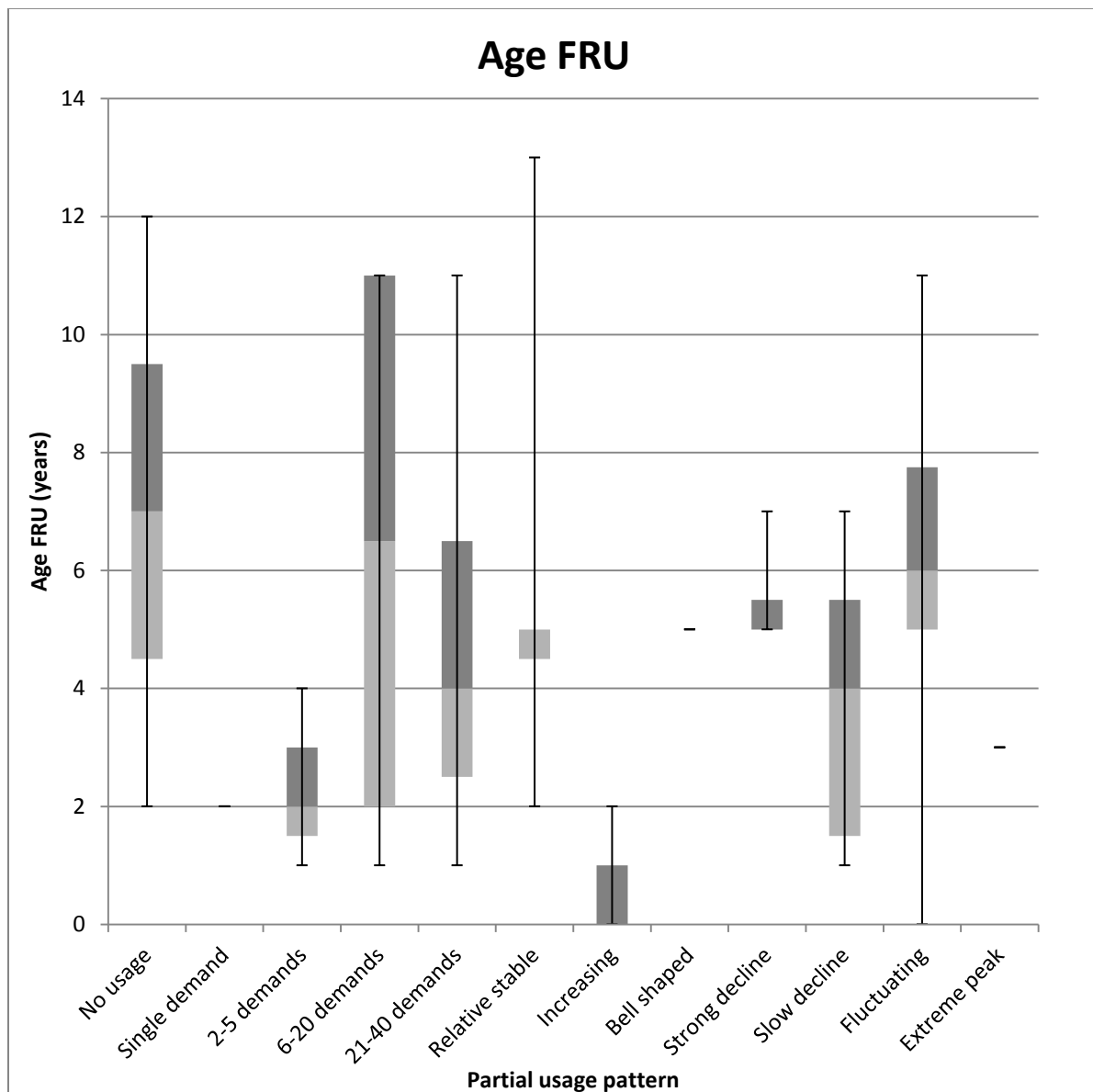


Figure H-2: Box plot Age FRU test selection

I. Description commodities CBLF

In section 6.2.1 the commodities applied in CBLF are divided over 3 groups, which indicate to which extent the tested forecasting methods lead to good results. These commodities differ from the commodities which are described in appendix E, because the CBLF commodities are grouped based on expected resemblance. Therefore there are more commodities in CBLF, consisting of a smaller amount of FRUs. A short explanation of the commodities will be given here.

AC-DC converter

An AC-DC converter is an adapter which transfers alternating current to direct current. Alternating current is the type of current delivered via the national power supplies, while direct current is the type of current required by devices to operate. The AC-DC converter is the device that forms the connection between the power plug and the machine.

Active backplanes

An active backplane is a panel consisting of various connectors, mini computers and printed circuit boards. Chips are added to buffer signals when required. There are no storage or processing devices directly mounted on the backplane. When storage or processing devices are required, they are added via plug-ins. An example of an active backplane can be found in Figure I-1.

Bezels, doors and fillers

The commodity Bezels, doors and fillers consists of a number of parts that are important for the outside of the machine, like the plastic or metal cover or the doors of a large machine.

CD/DVD drives

CD and DVD drives are the devices that can read the information on a CD or DVD. They can e.g. be found in laptops, but also in storage devices.

Custom function cards

Custom function cards are plug-in cards, that can be used for various reasons. In most cases, a set of functions is preprogrammed on the card, and inserting the card in the machine will create the opportunity to use these functions.

Display/video cards

This commodity consists out of two subgroups, being displays and video cards. A display is the interface which displays the graphical output of the machine, while a video card is the device that converts the output of the machine into images that will be shown on the display. In most cases, a video card is a plug in card, which can be added to the backplane.

Ethernet adapters

An Ethernet adapter is plug-in device that establishes the connection between the device the Ethernet adapter is in and the network of devices the device should operate in. In case of a laptop,

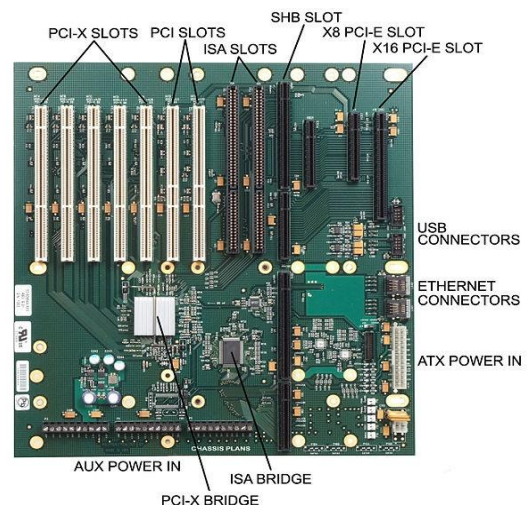


Figure I-1: Active backplane

the Ethernet adapter establishes the connection between the laptop and the LAN at e.g. a company, such that the laptop is part of the LAN environment.

Fans

A fan is a mechanical device, most commonly used to cool the machine it is mounted in. A fan can be found in many devices, like computers and storage devices. In all cases, the machines create heat during the time they are operational, which could cause damage. The fan makes sure a large part of the heat is blown out of the machine, to keep it operational.

Hard disk drives

A hard disk drive is used for storing data, and can be found in e.g. computers. The data on the hard disk drive consists out of data required for processes, but also out of user created files. There are different types of hard disks, related to the methods used to store data, and differences in storage size. All existing hard disk drive types are combined in one commodity.

I/O riser cards



Figure I-2: I/O riser card

An I/O riser card, of which an image can be found in Figure I-2, is a circuit board that is used for input/output actions for plug-in cards. The input consists out of signals from the plug-in card, which are transferred to the aspect of the machine for which the information is meant.

Keyboards

Keyboards are external or partially integrated devices that provide the opportunity for humans to communicate with a machine. Keyboards are most commonly known from computers, but can also be incorporated into mainframes and storage devices, to be able to operate these machines.

Mech ass and sub ass

The commodity mechanical assemblies and sub assemblies consists of all kinds of assemblies that can be found in the devices of IBM. An example of a mechanical assembly is a fan rack, which is a rack in which multiple fans are incorporated, used to cool large machines. An example of a sub assembly could be a light path, which is a display panel consisting out of a number of LED lights, all used to indicate the status of aspects of the machine. Examples of the statuses that could be communicated via a light path are power on, low battery or system issues. An example of a light path for a storage unit can be found in Figure I-3.



Figure I-3: Light path

Memory subs

The commodity Memory subs consists of devices that provide memory, which the machine can use to store information while executing operations. Different sizes of memory devices are provided, but all are combined in this commodity.

Micro assemblies

The commodity micro assemblies covers the group of processors IBM offers. A processor is the device that executes the instructions of computer program, from receiving information to executing

calculations. The processors can differ in the speed with which they can execute operations, but also in the size of the processor itself, depending on the amount of space available in the machine.

Rechargeable battery

Rechargeable batteries are, as the name suggests, batteries that are rechargeable. The commodity includes batteries for e.g. laptops, which are considerably large, but also small button batteries that are used to store small bits of information, like date and time.

SCSI adapter

SCSI is a technology used to connect devices within a machine, like e.g. hard disk drives and tape drives. It is mostly used for machines that perform extensive processing. An SCSI adapter is the device that connects one device to one or more SCSI connected devices, see Figure I-4. As a result, data can be transferred between the devices.



Figure I-4: SCSI adapter

System planar

A system planar is also known as a motherboard, and is a panel on which e.g. the CPU and memory are placed. The difference between an active backplane and a system planar is, that the system planar has storage capacity.

Tape drives

Tape drives are storage devices that use tape to store data on. Tape drives can be used to store large amounts of data fast, at relatively low cost. Retrieving data from tape drives is possible, but is slow, because the tape has to be unrolled before the storage location can be accessed.

J. Standard distributions

One of the improvement options we considered with respect to CBLF is using standard distribution curves to forecast the usage pattern. We selected a set of standard distributions, of which a description will be given. The information of the Weibull, Gamma, Exponential, Lognormal and Beta distribution is based on (NIST/Sematech, 2012).

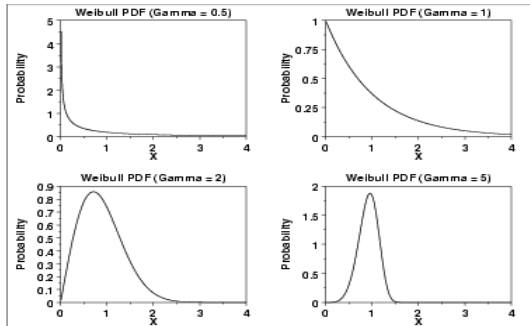


Figure J-1: Examples Weibull distribution

Gamma distribution

The gamma is a continuous distribution that can take on many different shapes, which are determined by the shape parameter, as shown in Figure J-2. The distribution is often used in computing times to the first failure of a system.

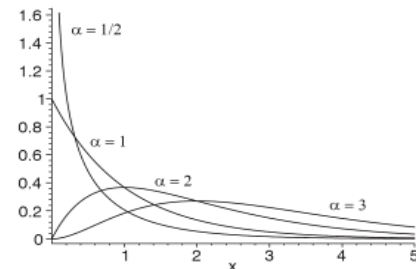


Figure J-2: Examples Gamma distribution

Exponential distribution

The exponential distribution is a distribution that is commonly used to model different quantities. Generally, its shape is equal to the shape of the Gamma distribution with $\lambda = 1$ and $\alpha = 1/2$, which can be found in Figure J-2, but an increasing shape is also possible.

Lognormal distribution

The lognormal distribution is a distribution in which the natural logarithm of the values X is normally distributed. The shapes the lognormal distribution can take on have either a short section with a strong increase in usage followed by an exponential distribution like shape, or the shape is similar to the shape of the exponential distribution from the start. As the Weibull distribution, the lognormal distribution is mostly used in reliability theory.

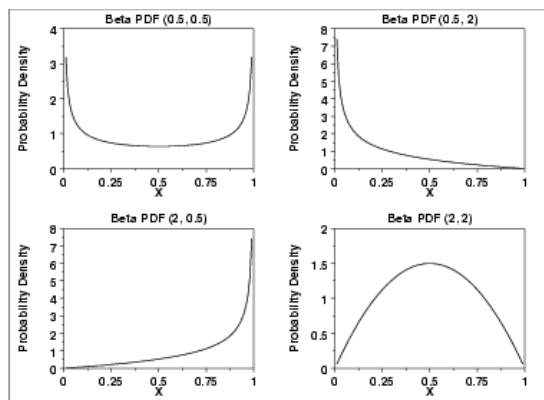


Figure J-3: Examples Beta distribution

Beta distribution

The Beta distribution is a distribution with a lower and upper bound, that indicate between which periods the Beta distribution can determine the shape of the curve. As the previous distributions, the Beta distribution also has a shape parameter, that determines the shape the curve actually takes on. Possible shapes the Beta distribution can take on are shown in Figure J-3 for the case the lower bound is 0 and the upper bound is 1. Similar shaped can be realized for situations with different lower and upper bounds.

Linear distribution

The linear distribution we apply is based on linear regression, which determines the straight line with the best possible fit through the historical data. Because a declining linear curve can lead to negative values, and negative usage is not possible, the selected curve value will always be the maximum of zero and the value based on the period and the linear regression parameters. This makes sure negative values cannot occur.

Constant distribution

The constant distribution we applied is a variation to the uniform distribution. The value of the distribution function is constant over the entire PLC, but does not necessarily has to be 1. It can take on the value that leads to the best result for the period of which historical usage is known. Because the forecast with the constant distribution is the same for all periods, we cannot calculate the value of the correlation coefficient R for this approach. Therefore we focus on minimizing the difference between the forecasted amount and the actual usage over the historical usage, based on the assumption that a correct forecast over historical usage is likely to result in a correct forecast over future period.

Croston's method

Croston's forecasting approach is based on separate forecasting of the periods in which usage occurs and the amount used at that point in time. The historical usage is the input required for the determination of the parameters of both the time between the usage occasions and the amount used.

Bass distribution

The Bass model (Bass, 2004) focuses on forecasting the future sales of a product. The forecast is created based on three parameters, being the market potential M, the coefficient of innovation P and the coefficient of imitation Q. The market potential indicates the total amount of products that is expected to be sold over the entire lifetime. The coefficient of innovation P indicates whether customers are likely to buy the product as soon as it gets on the market, while the coefficient of imitation Q focuses on the probability a customer buys the product when he or she sees many others already have the product. Based on the values of the parameters, different shapes can be created.